

Université Paris-Sud
UFR Scientifique d'Orsay

Thèse présentée pour obtenir le grade de
Docteur en Sciences de l'Université Paris XI

Spécialité :
Physique Théorique

présentée par
Thierry Mora

Géométrie et Inférence
dans l'optimisation et en théorie de
l'information

Soutenue le 24 septembre 2007
devant le jury composé de

Olivier Dubois	<i>Examineur</i>
Olivier Martin	<i>Examineur</i>
Marc Mézard	<i>Directeur de thèse</i>
Nicolas Sourlas	<i>Président</i>
Federico Ricci-Tersenghi	<i>Rapporteur</i>
Martin Weigt	<i>Rapporteur</i>

Remerciements

Je tiens tout d'abord à remercier mon directeur de thèse, Marc Mézard, dont l'attention bienveillante et les précieux conseils m'ont permis de mener à bien ce travail. Je voudrais aussi remercier Federico Ricci-Tersenghi et Martin Weigt d'avoir rempli le rôle de rapporteur, ainsi qu'Olivier Dubois, Olivier Martin et Nicolas Sourlas d'avoir accepté de figurer dans mon jury.

Je remercie également tous les chercheurs avec lesquels j'ai eu l'occasion de collaborer ou de discuter pendant ma thèse. Je suis notamment reconnaissant à Riccardo Zecchina de m'avoir invité à Trieste à plusieurs reprises, et de m'avoir initié à de nombreux problèmes passionnants. Ma gratitude va aussi à mes deux collègues, Olivier Rivoire et Lenka Zdeborova, dont j'ai beaucoup appris.

Le cadre de travail dont j'ai bénéficié a joué un rôle important dans l'élaboration de cette thèse. Je remercie donc les membres de mon laboratoire d'accueil, le LPTMS, en particulier son directeur Stéphane Ouvry, ses secrétaires Claudine Le Vaou et Martine Thouvenot, ainsi que ses ingénieurs système Olivier Brand-Foissac et Vincent Degat. Mes remerciements tout particuliers vont aux thésards du labo, parmi lesquels Yacine Ikhlef (pour ses bons mots) et Jérôme Roccia (pour son sens de l'à-propos), ainsi qu'à Benjamin Preciado et Michel Givort de l'IPN, qui m'ont accompagné moralement tout au long de ces trois années.

Table des matières

Remerciements	iii
Introduction	1
1 Approche physique de la théorie de l'information	9
1.1 Principes de la théorie de l'information	9
1.1.1 Information et entropie	9
1.1.2 Entropie physique	11
1.1.3 Limite thermodynamique	12
1.1.4 Exemples	16
1.2 Codage	19
1.2.1 Communication par un canal bruité	19
1.2.2 Codes aléatoires	22
1.2.3 Compression avec perte	27
1.2.4 Effets de taille finie	29
2 Approche physique de la complexité	33
2.1 Théorie classique de la complexité	33
2.1.1 Optimisation combinatoire	33
2.1.2 P vs NP	35
2.2 Complexité du pire et complexité typique	36
2.2.1 Motivation	36
2.2.2 Ensembles aléatoires et transitions de phase	37
2.2.3 Les problèmes réels sont-ils aléatoires ?	39
2.3 Diagramme de phases	40
2.3.1 Formulation physique	40
2.3.2 Fragmentation et condensation	42
2.3.3 Modèle à amas aléatoires	46
2.3.4 Ergodicité	48
3 Modèles graphiques	51
3.1 Graphes et hypergraphes	51

3.1.1	Graphes aléatoires	51
3.1.2	Coloriage	52
3.1.3	Graphes factoriels	53
3.2	Équations linéaires booléennes	55
3.2.1	Le problème XORSAT aléatoire	56
3.2.2	Utilisation pour la compression de données	58
3.2.3	Les codes linéaires dilués	60
3.3	Problèmes d'occupation	63
4	Passage de messages	67
4.1	Approximation des arbres	67
4.1.1	Chaîne d'Ising	68
4.1.2	Ramification de branches	70
4.1.3	Extension aux graphes dilués	72
4.1.4	Propagation des convictions	75
4.1.5	Statistique sur les instances	75
4.1.6	Stabilité et restructurabilité	76
4.2	Exemples	78
4.2.1	Décodage itératif	78
4.2.2	Énumération des A -parties d'un graphe factoriel	80
4.3	Calcul des corrélations	83
4.3.1	Propagation des susceptibilités	83
4.3.2	Application : modèles d'entropie maximale	85
5	Spectres de distance	91
5.1	Préliminaires : un peu de combinatoire	91
5.1.1	Le calcul recuit	92
5.1.2	Comparaison avec la moyenne gelée	94
5.1.3	Ensemble « lâche »	94
5.2	x -satisfaisabilité et fragmentation	96
5.2.1	x -satisfaisabilité dans k -XORSAT	98
5.2.2	x -satisfaisabilité dans k -SAT	100
5.2.3	L' x -satisfaisabilité dans le modèle à amas aléatoires	106
5.3	Distances et erreur dans les codes linéaires	108
5.3.1	Ensemble expurgé	109
5.3.2	Bornes d'union	109
6	Statistique des amas	115
6.1	Statistique des convictions	115
6.1.1	Une mesure sur les états	115
6.1.2	Propagation des sondages	116

6.1.3	Réduction à un état unique et condensation	119
6.1.4	Le seuil de satisfaisabilité	120
6.2	Modèles étendus	126
6.2.1	Fonction d'énumération du gel	127
6.2.2	Blanchissement	129
6.3	Retour sur les distances	130
6.3.1	Diamètre	130
6.3.2	Distances entre amas	133
Conclusion		139
Articles		145
	Clustering of Solutions in the Random Satisfiability Problem	147
	Pairs of SAT assignments in Random Boolean Formulæ	153
	Geometrical organization of solutions to random linear Boolean... .	179
	Error Exponents of Low-Density Parity-Check Codes...	203
	Statistical mechanics of error exponents for error-correcting codes . . .	211
Bibliographie		250

Introduction

Les problèmes d'optimisation combinatoire se manifestent sous des apparences variées dans de nombreux domaines de la science. Sous cette appellation se cache un principe très simple : étant donnée une fonction de coût définie sur un espace discret de variables, comment minimiser ce coût ? En physique par exemple, la minimisation de l'énergie réalise un principe d'optimisation, comme l'illustrent les problèmes de verre de spin, *cf.* figure 1. Outre les applications industrielles évidentes, dans la planification et dans l'organisation par exemple, les tâches d'optimisation sont aussi au cœur des méthodes d'inférence et d'apprentissage, où il s'agit de reconstruire une information incomplète à partir de règles établies au préalable, c'est-à-dire, de reconstruire le contenu le plus probable selon un principe de maximisation de la vraisemblance.

Au contraire des problèmes continus de basse dimension, pour lesquels il existe la méthode générique, et souvent efficace, de descente maximale selon la direction du gradient, les problèmes d'optimisation discrète souffrent en général du mal de la *frustration*. Dans le problème de verres de spins par exemple, évoqué figure 1, il peut arriver que certains groupes de spins, mis en présence de desiderata contradictoires induits par leurs voisins, demeurent dans l'indécision, faute de stratégie univoque telle que celle du gradient. Ainsi, les choix locaux échouent à assurer la minimisation *globale* de la fonction de coût. Ce phénomène de frustration explique et illustre la difficulté que représentent beaucoup de problèmes d'optimisation : ceux-ci peuvent parfois requérir

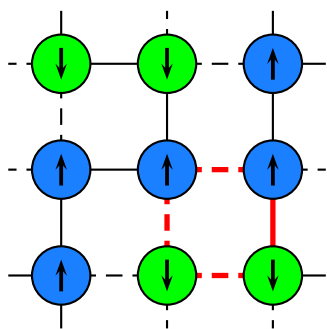


Fig. 1: Un problème de verre de spins. Sur chaque nœud repose une variable de spin, qui peut pointer dans deux directions opposées. Un lien plein entre deux nœuds signifie qu'il est plus favorable énergétiquement pour les deux spins de pointer dans la même direction (interaction ferromagnétique). Un lien en pointillés correspond à la situation inverse (interaction antiferromagnétique). La frustration est mise en évidence sur une boucle de liens contradictoires entre eux (traits rouge épais).

« Vous êtes le chef du protocole pour le bal de l'ambassade. Le prince héritier vous donne pour instruction soit d'inviter le Pérou, soit d'exclure le Qatar. La reine vous demande d'inviter le Qatar, la Roumanie ou les deux. Le roi, d'humeur rancunière, veut snobber la Roumanie, le Qatar ou les deux. Existe-t-il une liste d'invités qui satisfasse tous les caprices de la famille royale ? »

Ce problème de *satisfaisabilité*, tiré de l'article de vulgarisation de Bryan Hayes [Hay97], peut être formalisé par :

$$(p \text{ OU } \neg q) \text{ ET } (q \text{ OU } r) \text{ ET } (\neg r \text{ OU } \neg q)$$

où p , q et r sont les variables booléennes codant la présence des pays sur la liste d'invités.

Fig. 2: Le problème de satisfaisabilité. On cherche à trouver une solution à un ensemble de *clauses* sous forme conjonctive. Quand la formule et le nombre de variables deviennent grands, la question de la satisfaisabilité (*i. e.* de savoir si la frustration percole) devient difficile.

d'examiner un nombre prohibitif de configurations avant de trouver la solution optimale. La formalisation de cette difficulté a fait l'objet d'efforts importants de la part de la communauté des sciences informatiques, donnant naissance à la théorie de la complexité algorithmique. Elle s'applique en premier lieu aux problèmes de satisfaction de contraintes, une sous-classe de problèmes d'optimisation où l'on impose que le coût soit ramené en dessous d'une certaine valeur. Le problème de satisfaisabilité, présenté figure 2, en est l'exemple le plus illustre. Si l'on en croit la célèbre conjecture $P \neq NP$, encore indémontrée à ce jour, beaucoup de ces problèmes n'admettent une résolution qu'au terme d'un nombre d'opérations arithmétiques croissant *exponentiellement* avec la taille du problème, quelle que soit la procédure algorithmique retenue. Ces longs temps subis par les algorithmes de résolution ont depuis longtemps été rapprochés du phénomène de *trempe*, par lequel un matériau, soudainement refroidi à basse température, peine à se thermaliser et à réduire son énergie interne, par la faute d'une accumulation de frustrations locales dont la somme ne peut être surmontée que par un réarrangement à grande échelle du système. Une telle phase bloquée est qualifiée de « vitreuse ».

Cette analogie entre les problèmes d'optimisation et les verres a permis de tisser des liens fructueux, tant conceptuels que techniques, entre physique statistique et complexité algorithmique. L'un des enjeux principaux était, et demeure, de décider si le caractère intraitable de certains problèmes combinatoires admet une interprétation physiquement intelligible. Une voie prometteuse, qui a servi de fil conducteur à ce travail de thèse, a consisté à mettre en rapport la difficulté algorithmique avec les

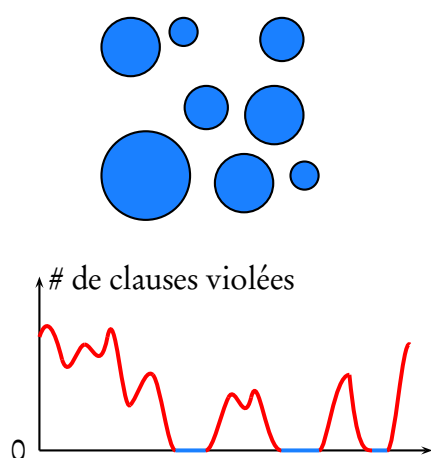


Fig. 3: Phénomène de fragmentation. L'espace des solutions d'un problème de satisfaction de contraintes, de très haute dimension, est ici schématiquement représenté en deux dimensions (en haut). Les amas forment une partition d'amas disjoints de solutions. Cette brisure spatiale est mise en rapport avec la brisure d'ergodicité observée dans les algorithmes de résolution. La représentation unidimensionnelle en fonction du nombre de clauses violées (en bas) met en évidence les hautes barrières qui séparent les amas. Ce trait est caractéristique des phases vitreuses.

propriétés purement géométriques de l'espace configurationnel : il a ainsi été conjecturé que l'espace des solutions d'un problème de satisfaction de contraintes pouvait souvent se fragmenter en une partition disjointe de composantes connexes, appelées amas, *cf.* figure 3. Cette conjecture fournit l'un des ingrédients les plus importants sous-tendant l'hypothèse de « brisure de symétrie de répliques » exploitée par les physiciens dans l'étude statistique des problèmes de satisfaction de contraintes, qui a notamment permis l'établissement de diagrammes de phases exacts et de stratégies algorithmiques novatrices. Basé sur des arguments heuristiques cohérents, cette hypothèse manquait toutefois de fondements mathématiques rigoureux. Par l'étude du spectre des distances, nous avons pu établir dans cette thèse la réalité du phénomène de fragmentation dans le très étudié problème de la satisfaisabilité.

Si ces phénomènes de fragmentation peuvent être rendus pour une part responsables des difficultés rencontrées dans la résolution des problèmes d'optimisation, ils peuvent aussi être mis à profit dans le contexte de la théorie de l'information : les « amas » étant identifiés aux messages possibles d'une source d'information, les propriétés de séparabilité peuvent servir à des fins de correction d'erreur. La figure 4 illustre schématiquement cette observation. On identifie des messages possibles (lettres, mots, chaînes de caractères, etc.) à certains points isolés, selon une cartographie prédéfinie et connue du destinataire. Sachant que le bruit peut, lors de la transmission, écarter le message de son point original, on a intuitivement intérêt à éloigner autant que possible les points les uns des autres. De fait, la notion de distance minimale est centrale en théorie de codage, où elle est intimement liée aux caractéristiques de discrimination dans le décodage par vraisemblance maximale.

La structure géométrique de l'espace des solutions d'un problème de satisfaction de contraintes et son influence sur les propriétés d'inférence sont les principaux

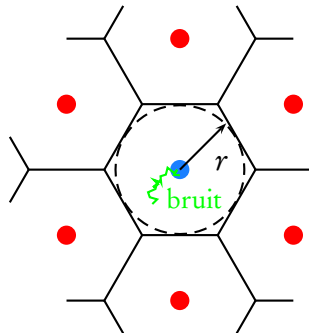


Fig. 4: Séparabilité dans les codes, illustrée ici sur l'espace réel à deux dimensions. Les « mots de code », c'est-à-dire les messages à transmettre, sont représentés par des points. Après la corruption par le canal de communication, modélisée ici par une marche aléatoire, le mot original en bleu pourra être récupéré à condition que la marche reste à l'intérieur de sa cellule de Voronoï. Pour les vrais codes, l'espace approprié est celui des longues chaînes de variables binaires, qui est de très haute dimension. L'importance de la *direction* prise par le bruit y est plus grande : en particulier, la sphère de sécurité de rayon r devient une très mauvaise approximation de la cellule de Voronoï.

thèmes abordés dans cette thèse. Ces notions sont exploitées dans deux domaines apparemment distincts : d'une part la théorie de la complexité algorithmique dite « typique », par opposition à celle du pire des cas, et d'autre part la théorie de l'information. Des ponts sont jetés entre ces domaines, sous l'égide de la physique statistique, dont la référence porte plus sur un corpus de notions et de méthodes que sur la réalité des phénomènes naturels. Néanmoins, quand cela s'avèrera utile, les origines des concepts physiques seront rappelées au fur et à mesure de leur introduction.

Contenu de la thèse

Le travail de thèse proprement dit s'articule autour de deux thématiques apparemment distinctes. D'une part, l'étude mathématique des propriétés de distance dans les problèmes de satisfaction de contraintes ; d'autre part, l'estimation de la probabilité d'erreur dans une classe particulière de codes de corrections d'erreur. L'exposé s'efforce, autant que possible, de rassembler dans un cadre commun inspiré de la physique statistique les concepts-clés intervenant au sein de chacune de ces thématiques. La présentation de cette thèse ne vise pas à la rigueur mathématique, bien que certains des résultats exposés soient rigoureusement établis.

Dans un premier chapitre, nous introduisons les principes fondamentaux de la théorie de l'information. Outre leur utilité évidente pour les codes de correction d'erreur, ces principes sont rapprochés des fondements de la physique statistique d'équilibre, qui sert de cadre méthodologique général à la thèse. Dans un deuxième chapitre, les concepts et problématiques de la complexité algorithmique sont présentés. Les motivations théoriques et pratiques afférentes sont évoquées, et les résultats et conjectures les plus intéressants résumés, notamment le phénomène de fragmentation. Le troisième chapitre introduit les modèles graphiques, communs aux codes de correction

d'erreur et aux problèmes de satisfaction de contraintes. Y est entamée une discussion sur les liens entre géométrie et inférence, par le biais des algorithmes d'effeuillage. Le quatrième chapitre résume les techniques de passage de messages, qui occupent une place importante dans la thèse, et donne quelques exemples, notamment en théorie de l'information. La généralisation de ces techniques aux susceptibilités est également introduite et illustrée sur des exemples simples. Le cinquième chapitre aborde les calculs « recuits » de spectre de distances. Ces calculs, bien qu'approchés, permettent de dériver des bornes rigoureuses sur les distances extrêmes. Ils sont appliqués aux problèmes de satisfaction de contraintes — où ils servent à étudier le phénomène de fragmentation — et de codage — où la relation entre séparabilité et distance est précisée. Enfin, le sixième chapitre présente les techniques de passage de messages en présence d'une phase fragmentée. Certains résultats classiques sur les seuils de satisfaisabilité y sont dérivés à titre illustratif. Ces techniques sont également mises à profit afin d'accéder aux propriétés fines des amas, au premier rang desquelles les propriétés de gel et de distances.

Ces chapitres visent à mettre en contexte les travaux originaux de cette thèse, rassemblés dans la seconde moitié du texte, et dont le contenu est ici brièvement résumé.

Fragmentation et x -satisfaisabilité. Les articles [MMZ05a, MMZ05b], écrits en collaboration avec Marc Mézard et Riccardo Zecchina à destination des physiciens et des mathématiciens respectivement, établissent rigoureusement l'existence d'une phase fragmentée dans le problème k -SAT. Ce résultat confirme une conjecture auparavant proposée par la communauté de la physique statistique, en rapport avec la nature supposément vitreuse de l'espace des solutions dans la phase « difficile » des problèmes de satisfaction de contraintes. La preuve s'appuie sur la notion de x -satisfaisabilité, forgée pour l'occasion, et équivalente à celle de spectre de distances. Elle consiste à établir des bornes sur ce spectre, et à en déduire un critère suffisant pour la fragmentation. Le raisonnement développé dans ces articles est détaillé au paragraphe 5.2.2, où il est replacé dans le contexte général des spectres de distances. Dans l'article [MM06b], écrit en collaboration avec Marc Mézard, cette même x -satisfaisabilité est étudiée dans le problème k -XORSAT, où le spectre de distances est calculé exactement à l'aide des techniques de passage de messages développées au chapitre 6. Ce calcul est repris au paragraphe 6.3.

Probabilité d'échec dans la correction d'erreur. Développé dans les articles [MR06a, MR06b] écrits en collaboration avec Olivier Rivoire, ce travail met au point une technique générale pour calculer la probabilité d'erreur dans le décodage optimal des codes de correction d'erreur « LDPC » (codes linéaires booléens et dilués). Un formalisme thermodynamique, partiellement exposé au paragraphe 1.2.4, y est introduit. La méthode employée est basée sur la *méthode de la cavité avec grandes déviations* [Riv05], qui est une extension de la méthode de la cavité exposée au chapitre 4. Il s'agit en fait d'une version « grandes déviations » du calcul de cavité effectué au

paragraphe 4.2.1. Des transitions de phases « atypiques » sont mises en évidence, qui trouvent une interprétation en termes de spectres de distances, en rapport avec le calcul du paragraphe 5.3. Les calculs principaux de ces deux articles ne sont pas détaillés dans le présent exposé, car ils auraient nécessité un exposé complet de la méthode de la cavité avec grandes déviations. Nous avons plutôt préféré insister sur l'interprétation géométrique de l'erreur, en concordance avec le thème unificateur de la thèse.

En plus de ces articles, la thèse contient quelques éléments originaux non publiés. Tout d'abord, la propagation des susceptibilités, et son application au problème de la machine de Boltzmann, est présentée au paragraphe 4.3. L'affinage de la technique numérique correspondante en vue d'une application à des problèmes concrets d'inférence, issus de la biologie ou d'ailleurs, est l'objet d'un travail en cours. Ensuite, un modèle « à amas aléatoires » est introduit et analysé de manière extensive (§2.3.3 et §5.2.3). En dépit de son absence d'intérêt en tant que tel, ce modèle permet de mettre en lumière de nombreux concepts et calculs présentés dans d'autres problèmes plus intéressants, notamment les liens entre ergodicité, fragmentation et gel. Il généralise le modèle à codes aléatoires (similaire au modèle à énergies aléatoires de Derrida) en en proposant une version « floue ». Enfin, les calculs trempé (§4.2.2) et recuit (§5.1) des fonctions d'énumération des poids du problème général d'occupation (défini au §3.3), constituent une extension de plusieurs résultats obtenus auparavant dans des contextes particuliers. L'estimation trempée de la fonction d'énumération des sous-parties d'arrêt constitue un exemple d'application originale de ce calcul.

Notations

Les principales notations et abréviations sont répertoriées ici.

$\doteq$	« par définition »
$\log, \ln$	les logarithmes en base 2 et e
$a_N \sim b_N$	$\lim_{N \rightarrow \infty} a_N / b_N = 1$
$a_N \asymp b_N$	$\log a_N \sim \log b_N$
$H(x)$	l'entropie binaire de paramètre x : $-x \log(x) - (1-x) \log(1-x)$
$D(x y)$	divergence de Kullback-Leibler : $x \log(x/y) + (1-x) \log[(1-x)/(1-y)]$
$H(p_N)$	entropie de la distribution p_N : $-\sum_{\sigma} p_N(\sigma) \log p_N(\sigma)$
$\mathbb{I}(A), \mathbb{P}(A)$	la fonction indicatrice et la probabilité d'un événement
$\mathbb{E}(X)$	l'espérance d'une variable aléatoire X
$ A $	le cardinal d'un ensemble A
p.s.	presque sûrement, <i>i.e.</i> avec probabilité tendant vers 1
$\lceil x \rceil$	la valeur entière « plafond » d'un nombre réel x
$\lfloor x \rfloor$	la valeur entière « plancher » d'un nombre réel x
$\delta_{a,b}, \delta(a,b)$	la fonction de Dirac discrète
$\delta(x)$	la fonction de Dirac continue
$\ \sigma\ $	la norme d'un vecteur booléen $\ \sigma\ $, égale $\sum_i \sigma_i $
$\partial a, \partial i$	l'ensemble des voisins du facteur a , ou de la variable i
σ_a	$(\sigma_i)_{i \in \partial a}$
$\sigma_{a \setminus i}$	$(\sigma_j)_{j \in \partial a \setminus i}$

Les notations suivantes se réfèrent la plupart du temps à :

N, M	Le nombre de variables, et le nombre de clauses ou de facteurs
α	la densité de clauses ou de facteurs M/N
σ	une configuration appartenant à $\mathcal{X}^N$, $\{0, 1\}^N$ ou $\{1, \dots, q\}^N$
$p(\sigma), E(\sigma)$	la probabilité d'une configuration, et son énergie
$\chi_a(\sigma_a)$	le poids d'un facteur a
β, m	température inverse et température inverse interne
$Z(\beta), \mathcal{Z}(m)$	fonction de partition à un état unique, et à états multiples
S, E, F	entropie, énergie, énergie libre
s, e, f	$S/N, E/N, F/N$
$\phi(\beta), \psi(m)$	potentiels à un état unique, et à états multiples

Quand un concept ou un objet originellement baptisé en anglais n'admet pas de traduction standard, la traduction française que nous proposons est :

cluster	amas
clustering	fragmentation
low-density parity-check codes	codes linéaires dilués
stopping set	sous-partie d'arrêt
belief propagation	propagation des convictions
survey propagation	propagation des sondages
warning propagation	propagation des avertissements
quenched	gelé
annealed	recuit

Chapitre 1

Approche physique de la théorie de l'information

Dans ce chapitre sont introduits les concepts et résultats importants de la théorie de l'information. L'approche adoptée s'inspire délibérément de la physique statistique, et les liens entre ces deux branches de la science sont soulignés.

1.1 Principes de la théorie de l'information

1.1.1 Information et entropie

Les fondements de la théorie de l'information ont été établis par Shannon dans son article pionnier de 1948 [Sha48], où la notion d'entropie fut introduite comme la mesure d'information d'une source de messages aléatoires. Supposons qu'une source discrète produise une chaîne de N lettres σ_i , $i = 1, \dots, N$, appartenant à un alphabet $\mathcal{X}$. Cette source est modélisée par une loi de probabilité sur les messages réalisés, notée $p_N(\sigma)$. Un message particulier σ étant produit, quelle quantité d'« information » contient-il ? L'idée de Shannon fut d'assimiler la notion intuitive d'information à une mesure de la « surprise » : plus un message est improbable, autrement dit moins son occurrence est prévisible, plus l'information qu'il apporte est importante. Pour des raisons que nous justifions ci-après, le logarithme offre une mesure appropriée de cette surprise : ainsi le contenu d'information associé à l'occurrence de σ est quantifié par

$$\log \frac{1}{p_N(\sigma)}. \quad (1.1)$$

Si, comme ce sera le cas tout au long de cette thèse, le logarithme s'exprime en base 2, cette quantité s'exprime en *bits* d'information.

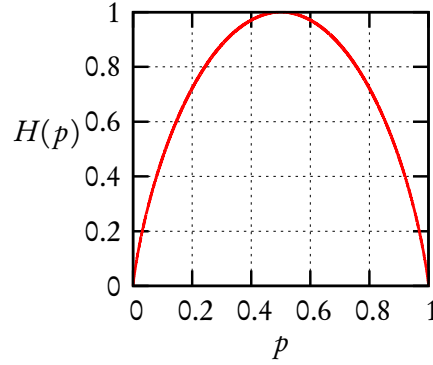


Fig. 1.1: Entropie binaire d'un processus de Bernoulli de paramètre p .

Afin de justifier le choix du logarithme, examinons le cas simple d'une pièce biaisée, qui prend la valeur « face » avec probabilité p , et « pile » avec probabilité $1 - p$. Si $p < 1/2$, les occurrences de « face » sont moins typiques, et contiennent donc plus d'information que celles de « pile ». Dans le cas extrême où $p = 0$, la pièce tombera toujours sur pile : cet événement étant parfaitement prévisible, aucune information supplémentaire n'est apportée. Au contraire, si la pièce n'est pas biaisée ($p = 1/2$), chaque occurrence apporte une information de un bit. En répétant N lancers de pièces non biaisées, nous produisons une séquence de piles et de faces, dont chacune a probabilité 2^{-N} . La quantité d'information est alors de $\log(1/2^{-N}) = N$ bits, ainsi que le suggère l'intuition.

L'entropie de la source est définie comme la moyenne de la mesure d'information sur les séquences :

$$H(p_N) = \mathbb{E} \log \frac{1}{p_N(\sigma)} = - \sum_{\sigma} p_N(\sigma) \log p_N(\sigma) \geq 0. \quad (1.2)$$

Reprenons l'exemple du lancer de pièce, aussi appelé processus de Bernoulli, afin de mettre en évidence les propriétés importantes de cette fonction. L'entropie vaut dans ce cas (cf. figure 1.1) :

$$H(p) = -p \log p - (1 - p) \log(1 - p) \quad (1.3)$$

- Cette fonction est maximale pour $p = 1/2$ où elle vaut 1 bit. De manière générale, l'entropie est maximale quand tous les messages sont équiprobables. En présence de Ω messages possibles et équiprobables, on a :

$$H = \log \Omega. \quad (1.4)$$

- $H(p)$ s'annule en $p = 0$ et en $p = 1$, illustrant la propriété qu'une source déterministe ne produit pas d'information.

- Quand N pièces sont lancées la probabilité de chaque séquence vaut :

$$p_N(\sigma) = \binom{N}{d} p^d (1-p)^{N-d}, \quad (1.5)$$

où d est le nombre d'occurrences de l'événement « face ». Le calcul de l'entropie de cette mesure donne $NH(p)$, conformément à un principe d'additivité. De manière plus générale, si on joint deux mesures indépendantes $p_N(\sigma)$ et $q_M(\tau)$, l'entropie du produit de ces deux mesures égale la somme des entropies de chacune :

$$H(p_N \otimes q_M) = H(p_N) + H(q_M). \quad (1.6)$$

Ces propriétés font de l'entropie une bonne candidate pour la mesure de l'information, et on peut même montrer qu'elle est la seule fonction à remplir ces conditions, à une constante multiplicative près.

1.1.2 Entropie physique

Bien avant que les bases de la théorie de l'information ne fussent jetées, l'entropie jouait déjà un rôle important en physique. En fait, l'introduction du concept d'entropie thermodynamique par Clausius en 1865 précède de presque un siècle la définition de Shannon. Plus tard, Boltzmann fut le premier à proposer une relation entre probabilité et entropie avec sa célèbre formule $S = k \log \Omega$ qu'il fit graver sur sa tombe. Cette formule est valable dans l'ensemble *microcanonique*, où l'on suppose que toutes les configurations d'une certaine énergie sont équiprobables : c'est le postulat de « désordre maximal ». L'entropie quantifie alors le nombre de configurations d'énergie E donnée :

$$\Omega(E) = 2^{S(E)} \quad (1.7)$$

(la constante multiplicative de Boltzmann k est ici fixée à 1).

Supposons que l'état d'un système physique soit décrit par un vecteur à valeurs discrètes $\sigma \in \mathcal{X}^N$ (position, spin, mode quantique, etc.). Quand le système est mis à l'équilibre avec un thermostat de température $T = \beta^{-1}$, la probabilité d'observer une configuration σ est donnée, sous l'hypothèse ergodique, par la loi de Boltzmann :

$$p_N(\sigma, \beta) = \frac{1}{Z(\beta)} 2^{-\beta E(\sigma)}, \quad (1.8)$$

où $E(\sigma)$ dénote l'énergie de la configuration σ , et $Z(\beta)$ est une constante de renormalisation, appelée *fonction de partition* :

$$Z(\beta) = \sum_{\sigma} 2^{-\beta E(\sigma)}, \quad (1.9)$$

reliée à l'énergie libre par $F(\beta) = \beta^{-1} \log Z(\beta)$. L'entropie de Gibbs est définie de la même manière que celle de Shannon :

$$H_N(\beta) = - \sum_{\sigma} p_N(\sigma, \beta) \log p_N(\sigma, \beta) = \beta \mathbb{E}[E(\sigma)] - \beta F(\beta). \quad (1.10)$$

La loi de Boltzmann sur les configurations définit l'ensemble *canonique*.

Ainsi que le suggèrent nos notations, cet ensemble est formellement identique au cas de la source aléatoire étudié au paragraphe précédent. Afin de mieux exploiter cette analogie, on peut associer à chaque message σ produit par une source une « énergie », ou log-vraisemblance :

$$E(\sigma) \doteq -\log p_N(\sigma), \quad (1.11)$$

et nous généralisons la mesure de probabilité par l'introduction d'une « température » fictive :

$$\begin{aligned} p_N(\sigma, \beta) &\propto p_N(\sigma)^\beta \\ &= \frac{1}{Z(\beta)} 2^{-\beta E(\sigma)} \end{aligned} \quad (1.12)$$

Cette généralisation est bien entendu formelle, et on se ramènera à $\beta = 1$ pour le cas réel.

Ainsi, le modèle de la source d'information s'inscrit naturellement dans le cadre de l'ensemble canonique.

1.1.3 Limite thermodynamique

Bien que les ensembles microcanoniques et canoniques reposent sur des postulats bien distincts, on peut montrer que ces deux niveaux de description sont équivalents dans la limite thermodynamique ($N \rightarrow \infty$). L'interprétation de cette limite diffère suivant que l'on se place du point de vue de la théorie de l'information ou de la physique statistique. En pratique, une source d'information est modélisée par une séquence aléatoires de lettres $\{\sigma_i\}_{i \geq 1}$. Cette séquence doit être suffisamment régulière pour que la loi de probabilité marginale $p_N(\sigma)$ des N premières lettres remplisse des conditions d'automoyennage. En particulier, la propriété d'additivité de l'entropie nous autorise à introduire un *taux d'entropie* :

$$h(\beta) = \lim_{N \rightarrow \infty} \frac{H_N(\beta)}{N}, \quad (1.13)$$

que nous supposons bien défini, et qui correspond à la quantité moyenne d'information par lettre. Du côté de la physique, la limite thermodynamique est généralement

justifiée par le grand nombre d'unités élémentaires (particules, spins, etc.) qui caractérisent les systèmes physiques. Là encore, on supposera l'existence d'une entropie de Gibbs par particule $h(\beta)$. Parallèlement, dans le contexte microcanonique, l'extensivité de l'entropie de Boltzmann conduit à postuler l'existence d'une entropie réduite :

$$s(e) = \lim_{N \rightarrow \infty} \frac{S(Ne)}{N}. \quad (1.14)$$

Équivalence des ensembles

Dans la limite des grands N , l'ensemble canonique concentre sa mesure autour d'une minorité de configurations équiprobables, dont le nombre est décrit par l'entropie de Gibbs $h(\beta)$. L'énergie se concentre autour de sa valeur moyenne $\mathbb{E}(E)$, et toutes les quantités thermodynamiques peuvent être déduites de l'ensemble microcanonique à énergie $\mathbb{E}(E)$. Ce scénario, dont nous prouvons la validité ci-dessous, correspond à ce qu'on appelle l'équivalence des ensembles.

Considérons dans un premier temps le comportement de la fonction de partition quand N tend vers l'infini :

$$Z(\beta) = \sum_{\sigma} 2^{-\beta E(\sigma)} = \sum_E 2^{S(E) - \beta E} \sim N \int_{-\infty}^{+\infty} de \, 2^{N[s(e) - \beta e]}, \quad (1.15)$$

La méthode de Laplace nous fournit un équivalent de cette quantité :

$$Z(\beta) \sim \sqrt{2\pi N} \left(-\frac{\partial^2 s}{\partial e^2} \Big|_{e^*} \right)^{-1/2} 2^{N[s(e^*) - \beta e^*]} \quad (1.16)$$

où e^* vérifie une équation de col :

$$\frac{\partial}{\partial e} (s(e) - \beta e) = 0, \quad \text{soit encore} \quad \beta = \frac{\partial s}{\partial e} \Big|_{e^*}. \quad (1.17)$$

L'énergie libre réduite $f(\beta) = \lim_{N \rightarrow \infty} F(\beta)/N$ prend alors une forme familière :

$$f(\beta) = e^* - \frac{1}{\beta} s(e^*). \quad (1.18)$$

L'estimation de l'énergie moyenne se fait également par la méthode de Laplace :

$$\mathbb{E}(E) = \frac{1}{Z(\beta)} \sum_{\sigma} E(\sigma) 2^{-\beta E(\sigma)} \sim N e^*, \quad (1.19)$$

ce qui entraîne l'égalité entre l'entropie de Gibbs (1.10) et celle de Boltzmann :

$$h(\beta) = s(e^*). \quad (1.20)$$

La méthode de Laplace repose de manière essentielle sur le fait que les sommes considérées sont exponentiellement dominées par le maximum de la fonction $s(e) - \beta e$, qui gouverne l'équilibre entre entropie et énergie. En pratique, cela signifie que presque toutes les configurations tirées au hasard réalisent cet équilibre. La démonstration de ce résultat requiert une inspection détaillée de la preuve du théorème de Laplace. Soit $\epsilon > 0$, et $A(\epsilon)$ l'ensemble des configurations d'énergie réduite $e \in (e^* - \epsilon, e^* + \epsilon)$. La somme des poids des configurations n'appartenant pas à $A(\epsilon)$ est majorée par :

$$\sum_{\sigma \notin A(\epsilon)} 2^{-\beta E(\sigma)} \sim N \left(\int_{e^* - \epsilon}^{e^*} + \int_{e^*}^{e^* + \epsilon} \right) de \, 2^{N[s(e) - \beta e]} \leq CN 2^{N[s(e^*) - \beta e^* - c\epsilon^2]}, \quad (1.21)$$

où C et c sont des constantes indépendantes de ϵ . À l'opposé, le poids cumulé de $A(\epsilon)$ est minoré par $2^{N[s(e^*) - \beta e^*]}$. Ainsi on a :

$$\frac{\mathbb{P}[\sigma \notin A(\epsilon)]}{\mathbb{P}[\sigma \in A(\epsilon)]} \leq CN 2^{-c\epsilon^2 N} \rightarrow 0, \quad (1.22)$$

Cela prouve que les configurations d'énergie arbitrairement proche de e^* dominent la mesure de Boltzmann. En outre, l'énergie donnant par définition une mesure de la vraisemblance des configurations, cf. (1.8), il en résulte que les configurations de $A(\epsilon)$ deviennent équiprobables quand $\epsilon \rightarrow 0$.

Cette observation entraîne une conséquence importante en théorie de l'information : la concentration de la mesure sur une minorité de messages implique la possibilité de *compresser* la source, c'est-à-dire de la représenter par un nombre réduit de lettres. En effet, avec une probabilité tendant vers 1, les messages produits appartiennent à $A(\epsilon)$, dont le cardinal :

$$2^{Ns(e^*)} \leq |A(\epsilon)| \leq N 2^{Ns(e^* + \epsilon)} \quad (1.23)$$

est gouverné par l'entropie de Gibbs $h(\beta) = s(e^*)$ avec une précision arbitraire.

Un procédé simple de compression consiste alors à numérotter ces messages de 1 à $|A(\epsilon)|$, et à leur associer un mot de $\lceil Ns(e^* + \epsilon) \rceil$ variables binaires. Les messages improbables $\notin A(\epsilon)$, dont le nombre est au plus $|\mathcal{X}|^N$, sont quant à eux codés par un mot binaire de taille $\lceil N \log |\mathcal{X}| \rceil$. La taille moyenne $\mathbb{E}(L_N)$ des mots ainsi obtenus vérifie :

$$\lceil Ns(e^* + \epsilon) \rceil \leq \mathbb{E}(L_N) \leq \lceil Ns(e^* + \epsilon) \rceil + CN 2^{-c\epsilon^2 N} \lceil N \log |\mathcal{X}| \rceil \quad (1.24)$$

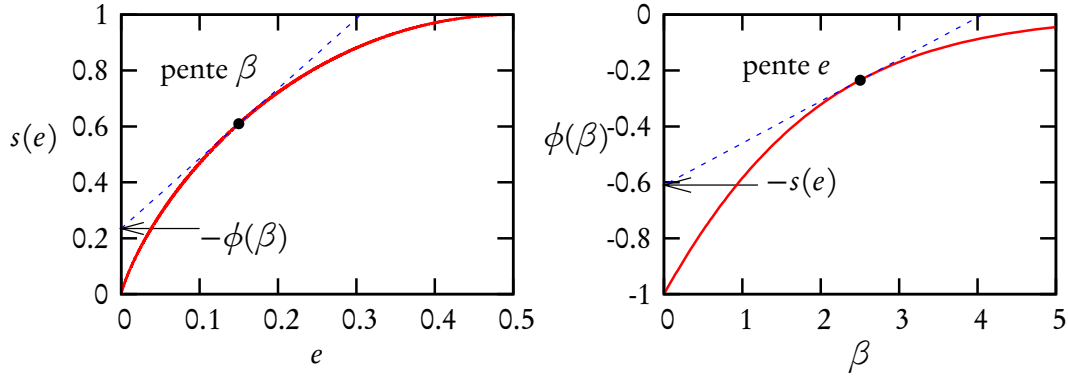


Fig. 1.2: Transformation de Legendre (panneau de gauche), et transformation inverse (panneau de droite). Le potentiel $\phi(\beta)$ est construit géométriquement en posant une droite de pente β sur la courbe $s(e)$. On a donc logiquement $\beta = s'(e)$. L'intersection de cette droite avec l'axe des ordonnées donne $-\phi(\beta)$ en vertu de (1.26). Réciproquement, $s(e)$ peut être construit à partir de $\phi(\beta)$ en utilisant (1.27).

En choisissant par exemple $\epsilon = N^{-1/3}$, on trouve :

$$\lim_{N \rightarrow \infty} \frac{\mathbb{E}(L_N)}{N} = s(e^*) = h(\beta). \quad (1.25)$$

Ce résultat important dû à Shannon, démontré ici dans le cadre de la physique statistique, prouve l'existence de codes de compression *optimaux*. La réciproque de ce théorème établit en effet l'impossibilité de compresser la source avec des mots binaires de taille moyenne $< Nh(\beta)$. Ce résultat est d'ailleurs plus facile à comprendre : si un tel procédé de compression existait, le taux d'entropie des mots compressés serait strictement inférieur à $h(\beta)$, ce qui mettrait en défaut sa capacité à représenter la source originale.

Une interprétation alternative de l'entropie est ainsi dégagée : d'après le théorème de Shannon, l'entropie mesure la « taille » de l'espace des messages typiques, ou, de manière équivalente, le nombre de lettres nécessaires pour en décrire les éléments.

Revenons un moment sur les conséquences de l'équivalence des ensembles, et notamment sur la formule reliant l'énergie libre à l'entropie :

$$\phi(\beta) = \min_e [\beta e - s(e)], \quad (1.26)$$

où $\phi(\beta) \doteq \beta f(\beta)$ est appelée *fonction de potentiel*.

On reconnaît dans cette relation une *transformation de Legendre* (voir figure 1.2), qui s'inverse comme suit :

$$s(e) = \min_{\beta} [\beta e - \phi(\beta)]. \quad (1.27)$$

Cette dernière relation peut d'ailleurs être dérivée directement à partir de la définition de l'entropie :

$$\begin{aligned} 2^{Ns(e)} &= \sum_{\sigma} \delta[E(\sigma), Ne] = \int_{-i\pi/\log 2}^{i\pi/\log 2} \frac{d\beta \log 2}{2\pi i} \sum_{\sigma} 2^{-\beta[E(\sigma) - Ne]} \\ &= \int_{-i\pi/\log 2}^{i\pi/\log 2} \frac{d\beta \log 2}{2\pi i} 2^{N[-\phi(\beta) + \beta e]} \asymp 2^{N \min_{\beta \in \mathbb{R}} [\beta e - \phi(\beta)]}, \end{aligned} \quad (1.28)$$

où l'on a utilisé la représentation intégrale de la fonction de Dirac discrète, et où la méthode du col dans le plan complexe a été employée pour obtenir le comportement asymptotique¹. Dans ce calcul simple, β joue le rôle d'un multiplicateur de Lagrange contraignant l'énergie, et sa valeur au col β^* est déterminée par cette dernière. Nous avons vu que, dans le passage du canonique au microcanonique, la température remplit une fonction semblable en prescrivant la valeur de l'énergie autour de laquelle la mesure se concentre. L'énergie et la température sont dites *conjuguées*.

Ainsi, les fonctions d'entropie $s(e)$ et de potentiel $\phi(\beta)$, associées respectivement aux ensembles microcanonique et canonique, offrent des descriptions équivalentes du système, et se déduisent l'une de l'autre par des transformations de Legendre (1.26) et (1.27). Dans la suite de cette thèse on mettra à profit cette équivalence, fondamentale en physique statistique, dans des contextes variés. En particulier, les irrégularités de la fonction d'entropie microcanonique fournissent des indications sur l'existence de transitions de phase en physique. Par exemple, la non-convexité peut se traduire par une transition de phase du premier ordre, auquel cas l'équilibre thermodynamique est réalisé par une construction de Maxwell. De la même manière, les discontinuités de l'entropie signalent souvent une transition de gel, ainsi que nous le verrons au paragraphe 1.2.2 dans le cas du modèle à codes aléatoires.

1.1.4 Exemples

La similitude formelle entre les systèmes physiques et les sources d'information peut être illustrée par quelques cas simples trouvant des interprétations dans les deux domaines. Considérons par exemple une série de processus de Bernoulli indépendants :

$$p_N(\sigma) = \prod_{i=1}^N p_i(\sigma_i) \quad (1.29)$$

où $p_i(0) = p_i$, et $p_i(1) = 1 - p_i$. Par la propriété d'additivité, l'entropie s'évalue simplement à $\sum_i H(p_i)$. La version physique de ce système est un modèle de spins indé-

¹La présence du « min » s'explique par le fait que le col dominant est ici un minimum dans la direction réelle de β .

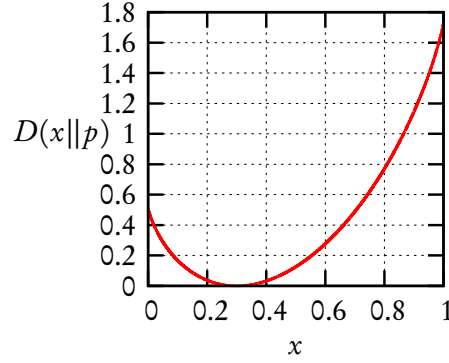


Fig. 1.3: Divergence de Kullback-Leibler $D(x||p=0,3)$ en fonction de x .

pendants $\sigma_i \in \{-1, 1\}$, d'énergie :

$$E(\sigma) = - \sum_i h_i \sigma_i. \quad (1.30)$$

relié au précédent par

$$\beta h_i = \frac{1}{2} \log \frac{p_i}{1-p_i}. \quad (1.31)$$

Quelles sont les configurations « typiques » de cet ensemble ? Le cas uniforme $p_i \equiv p$ permet de répondre simplement à cette question, sans pour autant dénaturer le comportement général. Le processus est alors équivalent à plusieurs lancers d'une pièce biaisée. Le comportement asymptotique de la loi de probabilité (1.5) est donné (en fonction du nombre d de « 1 ») par :

$$p_N(\sigma) = \binom{N}{d} p^d (1-p)^{N-d} \asymp 2^{N[H(x) + x \log p + (1-x) \log(1-p)]} \asymp 2^{-ND(x||p)}, \quad (1.32)$$

où la formule de Stirling $n! \sim \sqrt{2\pi n} n^n e^{-n}$ a été utilisée, et où $x = d/N$, la proportion de $\sigma_i = 1$, a été introduite. L'exponentielle est gouvernée par la *divergence de Kullback-Leibler* à deux éléments (cf. figure 1.3) :

$$D(x||p) = x \log \frac{x}{p} + (1-x) \log \frac{1-x}{1-p}. \quad (1.33)$$

Cette divergence est minimale pour $x = p$: les réalisations typiques du hasard sont donc celles où la proportion de « 1 » est proche de p . Le nombre de telles réalisations est asymptotiquement déterminé par l'entropie de Gibbs : $\binom{N}{pN} \asymp 2^{NH(p)}$.

Les sources naturelles d'information, comme par exemple les langues naturelles, sont loin d'être des suites de lettres tirées indépendamment. Bien au contraire, l'histoire passée de la source influence fortement son comportement futur. Le besoin

d'une description un peu plus réaliste des sources d'information conduit à modéliser la source par une chaîne de Markov :

$$p_N(\sigma) = p_1(\sigma_1)q_2(\sigma_2|\sigma_1)q_3(\sigma_3|\sigma_2)\cdots q_N(\sigma_N|\sigma_{N-1}). \quad (1.34)$$

Les marginales de la loi de probabilité généralisée $p_N(\sigma, \beta) = p_N(\sigma)^\beta / Z(\beta)$ s'estiment itérativement par la méthode des matrices de transfert :

$$p_i(\sigma_i, \beta) = \frac{1}{z_i(\beta)} \sum_{\sigma_{i-1}} p_{i-1}(\sigma_{i-1}, \beta) q_i^\beta(\sigma_i|\sigma_{i-1}), \quad (1.35)$$

où $z_i(\beta)$ est une constante de normalisation. L'énergie libre $F(\beta) = -\beta^{-1} \log Z(\beta)$ s'évalue à :

$$F(\beta) = -\beta^{-1} \sum_i^N \log z_i(\beta). \quad (1.36)$$

d'où l'on déduit l'entropie par $H_N(\beta) = \beta^2 \partial F / \partial \beta$.

Quand le système est uniforme, $q_i \equiv q$, les marginales p_i convergent vers $\tilde{p}$ quand i tend vers l'infini. Cette probabilité vérifie la relation d'auto-cohérence :

$$\tilde{p}(\sigma, \beta) = \frac{1}{z(\beta)} \sum_{\sigma'} \tilde{p}(\sigma', \beta) q^\beta(\sigma|\sigma') \quad (1.37)$$

Dans le cas réel $\beta = 1$, cela donne

$$\tilde{p}(0) \doteq \tilde{p}(0, \beta = 1) = \frac{q(0|1)}{q(1|0) + q(0|1)}. \quad (1.38)$$

De la même façon, la valeur asymptotique de $z_i(\beta)$ vérifie :

$$-\log z(\beta = 1) = 0, \quad (1.39)$$

$$-\frac{\partial}{\partial \beta} \log z(\beta = 1) = -\sum_{\sigma, \sigma'} \tilde{p}(\sigma') q(\sigma|\sigma') \log q(\sigma|\sigma'). \quad (1.40)$$

d'où l'on déduit :

$$\begin{aligned} h &= \lim_{N \rightarrow \infty} \frac{H(p_N)}{N} = \beta^2 \left. \frac{\partial [-\beta^{-1} \log z(\beta)]}{\partial \beta} \right|_{\beta=1} \\ &= -\sum_{\sigma, \sigma'} \tilde{p}(\sigma') q(\sigma|\sigma') \log q(\sigma|\sigma'). \end{aligned} \quad (1.41)$$

La version physique de ce problème est simplement le modèle d'Ising à une dimension. En effet, une fois effectuée la transformation des variables binaires en variables de spins, les probabilités de transition peuvent s'écrire sous la forme :

$$q_i(\sigma|\sigma') = 2^{a_i \sigma + b_i \sigma' + c_i \sigma \sigma' + d_i}, \quad (1.42)$$

où b_i et d_i assurent la normalisation. Cette forme conduit naturellement à définir l'énergie ainsi :

$$E(\sigma) = h_1 \sigma_1 - \sum_{i=2}^N (b_i \sigma_i + J_i \sigma_i \sigma_{i-1}) \quad (1.43)$$

avec :

$$h_i = a_i + b_{i+1} \quad \text{et} \quad J_i = c_i. \quad (1.44)$$

Nous reviendrons plus tard sur ce modèle simple (§4.1.1), et en particulier sur la transformation *inverse*, qui consiste à traduire une chaîne de spins d'Ising en une chaîne bayésienne de probabilités conditionnelles. Ce genre d'opération est en effet au cœur des techniques de passages de message, décrites au chapitre 4.

1.2 Codage

1.2.1 Communication par un canal bruité

Le choix de l'entropie comme mesure de l'information implique que les sources transmettent en général moins d'information que ne le permet leur alphabet. Cette « perte » d'information, ou *redondance*, est reflétée par la structure corrélative de la statistique des messages. À quoi cette redondance sert-elle ? Dans l'exemple des langages naturels, la suppression ou la corruption d'un faible nombre de lettres peut être corrigée par le lecteur, pour peu que celui-ci jouisse d'une connaissance suffisante de la langue, c'est-à-dire d'une estimation de la vraisemblance *a priori* de mots ou de groupes de mots. Lors d'un tel processus de correction d'erreur, la redondance est donc utilisée à bon escient pour reconstituer le message. Réciproquement, les sources sans redondance, où chaque séquence de lettres est également probable, ne survivent pas à l'épreuve du bruit : ne bénéficiant d'aucune connaissance *a priori* de la statistique des messages, le lecteur n'a pas les moyens d'inférer le message original en cas de corruption.

Le problème de la communication s'énonce comme suit. Mettons que l'on veuille transmettre des messages par un canal bruité. Comment y ajouter artificiellement de la redondance afin de compenser l'effet du bruit ?

À chaque message m original de taille L , on associe un *mot de code* $\sigma^0 = f(m)$ de taille $N > L$: cette opération d'*encodage* produit la redondance nécessaire pour lutter contre les effets du bruit. L'ensemble des mots de codes possibles définit le *livre de code*, dénoté par $\mathcal{C}$. Le nombre de bits d'information par lettre dans le mot de code définit le *taux* du code :

$$R = \frac{L}{N}. \quad (1.45)$$

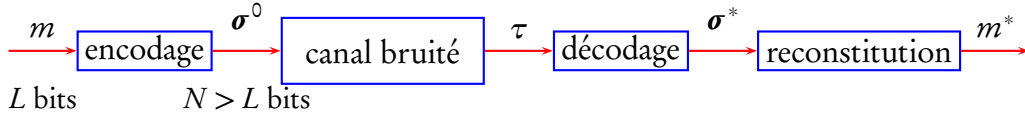


Fig. 1.4: Schéma de la communication sur un canal bruité.

Ce taux sera maintenu constant quand les longueurs N et L tendront vers l'infini.

Lors de son passage dans le canal, le mot de code est corrompu en une chaîne τ , selon une loi de probabilité conditionnelle $Q(\tau|\sigma^0)$, qui caractérise les propriétés du canal. En général, cette nouvelle chaîne ne fait pas partie du livre de code. Lors de la phase de décodage, on cherche le mot de code σ^* le plus vraisemblable, compte tenu de la chaîne reçue τ et des propriétés du canal. Enfin, le message $m^* = f^{-1}(\sigma^*)$ est reconstitué par inversion du livre de code. Le décodage sera un succès si $m^* = m$. La figure 1.4 résume le schéma ainsi proposé.

Le *théorème de Shannon sur les canaux discrets bruités* assure la possibilité de mettre en œuvre une communication sans erreur dans la limite des longs mots, à condition que le taux R du code ne dépasse pas la *capacité* $C(Q)$ du canal :

$$R < C(Q) \doteq \sup_{p_N} I(p_N, q_N), \quad (1.46)$$

où $p_N(\sigma)$ est une loi test des messages transmis, et $q_N(\tau) = \sum_{\sigma} p_N(\sigma) Q(\tau|\sigma)$ la loi consécutive des messages reçus. $I(p_N, q_N)$ désigne l'*information mutuelle* entre p_N et q_N , qui mesure leur dépendance statistique :

$$I(p_N, q_N) = \sum_{\sigma, \tau} p_N(\tau) Q(\tau|\sigma) \log \frac{p_N(\tau) Q(\tau|\sigma)}{p_N(\sigma) q_N(\tau)}. \quad (1.47)$$

Cette mesure s'interprète comme la réduction d'incertitude sur σ permise par la connaissance de τ . Dans le cas extrême où la relation entre ces deux mots est déterministe, la capacité vaut $C = \sup_{p_N} H(p_N) = 1$. À l'opposé, elle s'annule quand cette relation est complètement aléatoire, i.e. $Q(\tau|\sigma) = q_N(\tau)$.

Dans le cadre de cette thèse nous considérerons principalement deux types de canaux binaires ($\sigma_i \in \{0, 1\}$) sans mémoire, dont le comportement est illustré figure 1.5 :

- Le canal binaire symétrique (BSC). Chaque bit est inversé indépendamment avec probabilité ϵ :

$$Q(\tau|\sigma) = \prod_i \left[(1 - 2\epsilon) \delta_{\sigma_i, \tau_i} + \epsilon \right], \quad (1.48)$$

avec $\tau_i \in \{0, 1\}$. Sa capacité vaut $C_{\text{BSC}}(\epsilon) = 1 - H(\epsilon)$.

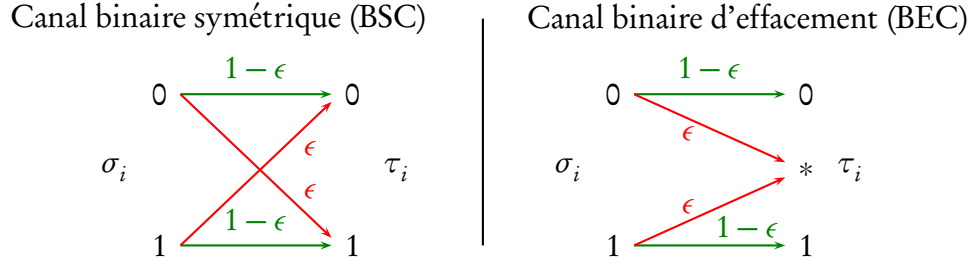


Fig. 1.5: Les deux canaux considérés dans cette thèse.

- Le canal binaire d'effacement (BEC). Chaque bit est effacé indépendamment avec probabilité ϵ :

$$Q(\tau|\sigma) = \prod_i \left[(1 - \epsilon) \delta_{\sigma_i, \tau_i} + \epsilon \delta_{\tau_i, *} \right], \quad (1.49)$$

où $*$ dénote un bit effacé, et $\tau_i \in \{0, 1, *\}$. La capacité de ce canal est $C_{\text{BEC}}(\epsilon) = 1 - \epsilon$.

Parmi les différentes tâches impliquées dans le schéma de la communication, le décodage est sans doute le plus difficile, bien que la génération des mots de codes puisse présenter des difficultés [Mac03]. La probabilité *a posteriori* du mot envoyé, conditionnellement au mot reçu, est exprimée par la formule de Bayes :

$$\mathbb{P}(\sigma|\tau) = \frac{Q(\tau|\sigma) \mathbb{I}(\sigma \in \mathcal{C})}{\sum_{\sigma'} Q(\tau|\sigma') \mathbb{I}(\sigma' \in \mathcal{C})}. \quad (1.50)$$

Il existe alors typiquement deux méthodes de décodage, suivant la définition d'erreur qu'on se donne :

- La vraisemblance maximale par mot : $\sigma^* = \operatorname{argmax}_{\sigma} \mathbb{P}(\sigma|\tau)$. Ce choix minimise la probabilité d'erreur de mot : $\mathbb{P}(\sigma^* \neq \sigma^0)$.
- La vraisemblance maximale par lettre : $\sigma_i^* = \operatorname{argmax}_{\sigma_i} \mathbb{P}(\sigma_i|\tau)$, qui maximise la marginale sur i . Ce choix optimise la probabilité d'erreur par lettre : $(1/N) \sum_i \mathbb{P}(\sigma_i^* \neq \sigma_i^0)$.

Ces deux méthodes peuvent être englobées dans un schéma plus général de décodage, reposant sur la définition d'une énergie :

$$E(\sigma) = -\log [Q(\tau|\sigma) \mathbb{I}(\sigma \in \mathcal{C})]. \quad (1.51)$$

Le décodage par vraisemblance de mot revient alors à chercher le *fondamental* de cet hamiltonien, i.e. la configuration d'énergie minimale. Dans le contexte canonique,

cela équivaut à prendre la limite de température nulle ($\beta \rightarrow \infty$). Le décodage par vraisemblance de lettres repose quant à lui sur l'estimation des marginales de la loi $\mathbb{P}(\sigma|\tau) \propto 2^{-\beta E(\sigma)}$, avec $\beta = 1$. L'interpolation entre ces deux cas s'obtient logiquement par le choix d'une température intermédiaire $1 \leq \beta \leq \infty$:

$$\sigma_i^* = \operatorname{argmax}_{\sigma_i} \sum_{\sigma_{\setminus i}} 2^{-\beta E(\sigma)}. \quad (1.52)$$

Remarquez que le mot reçu τ joue ici le rôle de variable « gelée », alors que σ regroupe les degrés de liberté du système. Le problème du décodage s'apparente ainsi à un système désordonné semblable aux verres de spins en physique.

Pour les canaux qui nous intéressent, l'énergie prend une forme simple. Avec le BSC,

$$E(\sigma) = \begin{cases} h \sum_i (1 - \delta_{\sigma_i, \tau_i}) & \text{si } \sigma \in \mathcal{C}, \\ +\infty & \text{sinon.} \end{cases} \quad (1.53)$$

où $h = \log[(1 - \epsilon)/\epsilon]$. L'énergie mesure la distance de Hamming au mot reçu, au facteur h près. La limite de température nulle correspond alors à la recherche du mot de code le plus proche de τ .

Le cas du BEC est encore plus simple : l'énergie y est uniforme pour tous les mots de code $\sigma \in \mathcal{C}$ compatibles avec le mot reçu (i.e. tels que pour tout i , $\tau_i \neq *$ implique $\sigma_i = \tau_i$), et elle vaut $+\infty$ dans le cas contraire. Du fait de cette dégénérescence, le décodage ne peut réussir que si l'ensemble des mots de codes compatibles avec le mot reçu est réduit à un singleton $\{\sigma^0\}$.

1.2.2 Codes aléatoires

La preuve originale de la partie existentielle du théorème de Shannon est notoirement non-constructive, en ce qu'elle s'appuie sur les propriétés moyennes d'un ensemble de codes aléatoires. Dans ce paragraphe nous développons cette preuve dans les cas particuliers du BSC et du BEC. Un code aléatoire binaire est construit en choisissant la fonction de codage f complètement au hasard, en associant à chacun des 2^L messages originaux un mot de code de longueur N tiré au hasard parmi les 2^N possibilités.

Outre leur utilisation dans la preuve de Shannon, les codes aléatoires possèdent des propriétés intéressantes en soi, et sont l'occasion d'introduire des concepts et de mettre en évidence des comportements que nous retrouverons par la suite dans les codes linéaires dilués. Comme nous le verrons dans le chapitre 4, les codes aléatoires peuvent d'ailleurs s'obtenir comme la limite de grande connectivité de tels codes. Parmi les notions centrales en théorie de la communication, les propriétés de distance, cruciales pour les questions de décodabilité, sont ici élucidées à l'aide de méthodes

combinatoires élémentaires. Les codes aléatoires affichent par ailleurs une phénoménologie très proche de celle du modèle à énergies aléatoires (*random energy model*, REM) introduit par Derrida et présenté comme le modèle archétypique de la transition vitreuse en physique statistique. Ils fournissent l'occasion d'illustrer un cas pathologique d'équivalence d'ensemble, où intervient une transition de condensation.

Sur le canal d'effacement

Quelle est la performance moyenne d'un code aléatoire quand il est utilisé sur le canal d'effacement ? Désignons par $E \subset \{1, \dots, N\}$ l'ensemble des bits effacés par le canal, et notons n le nombre de mots de codes distincts de σ^0 et compatibles avec le mot reçu, i.e. tels que $\sigma_i = \sigma_i^0$ pour tout $i \notin E$. Comme chaque mot de code autre que σ^0 est compatible avec probabilité $2^{|E|-N}$, n est une loi binomiale, avec notamment :

$$\mathbb{E}(n \text{ sachant } |E|) = (2^L - 1)2^{|E|-N}, \quad (1.54)$$

$$\text{Var}(n \text{ sachant } |E|) = (2^L - 1)2^{|E|-N} [1 - 2^{|E|-N}]. \quad (1.55)$$

Le nombre de bits effacés $|E|$ tombe dans l'intervalle $[N(\epsilon - \delta), N(\epsilon + \delta)]$ presque sûrement, pour tout $\delta > 0$. Supposons d'abord $L/N = R < 1 - \epsilon$. L'inégalité de Markov donne :

$$\mathbb{P}(n \geq 1) \leq \mathbb{E}(n) \leq \sum_{|E|} \mathbb{P}(|E|) 2^{N(R+|E|/N-1)} \leq 2^{N(R+\epsilon+\delta-1)}. \quad (1.56)$$

Cette probabilité tend vers 0 pour δ suffisamment petit et par conséquent, le décodage réussit presque sûrement.

Réciproquement, supposons maintenant $R > 1 - \epsilon$. Le théorème de Chebychev assure que n se concentre autour de sa valeur moyenne :

$$\mathbb{P} \left[\left| \frac{n}{\mathbb{E}(n)} - 1 \right| > C \text{ sachant } |E| \right] \leq \frac{\text{Var } n}{C^2 \mathbb{E}(n)^2} \leq C^{-2} 2^{-N(R+|E|/N-1)}. \quad (1.57)$$

Or, $R + |E|/N - 1 > (R - \epsilon - 1)/2 > 0$ presque sûrement, ce qui entraîne que le décodage échoue ($n > 0$) presque sûrement.

Ainsi, le code aléatoire subit une transition abrupte d'une phase décodable $\epsilon < \epsilon_c = 1 - R$ vers une phase indécodable $\epsilon > \epsilon_c$. Ce seuil de décodabilité réalise précisément la borne de Shannon dans le cas particulier du BEC :

$$R = C_{\text{BEC}}(\epsilon_c) = 1 - \epsilon_c. \quad (1.58)$$

Sur le canal symétrique

Le canal symétrique BSC requiert une analyse plus fine de l'espace des mots de code. Le rôle important que joue la distance dans l'estimation de la vraisemblance des

mots conduit à étudier le spectre des distances dans ces codes. Soit $\mathcal{C}' = \mathcal{C} \setminus \sigma_0$ un livre de code aléatoire duquel on a retiré le mot transmis. Le mot reçu τ est complètement aléatoire et indépendant de $\mathcal{C}'$. Désignons par n_w le nombre de mots de codes séparés de τ par une distance w :

$$n_w = \text{card}\{\sigma \in \mathcal{C}' \mid \|\sigma - \tau\| = w\} \quad (1.59)$$

où $\|\sigma - \tau\|$ dénote la distance de Hamming $\sum_i (1 - \delta_{\sigma_i, \tau_i})$. La fonction génératrice de n_w est connue sous le nom de *fonction d'énumération des poids* :

$$n(x) = \sum_{w \geq 0} n_w x^w. \quad (1.60)$$

Chaque mot de code σ est à distance w de τ avec probabilité :

$$p_w = \binom{N}{w} 2^{-N}. \quad (1.61)$$

Chacun de ces mots étant tirés indépendamment, n_w suit une loi binomiale :

$$\mathbb{E}(n_w) = (2^L - 1)p_w \asymp 2^{N[H(\omega) + R - 1]}, \quad (1.62)$$

$$\text{Var } n_w = (2^L - 1)p_w(1 - p_w) \asymp 2^{N[H(\omega) + R - 1]}. \quad (1.63)$$

où on a supposé que la distance w croît proportionnellement à N dans la limite asymptotique : $w = \omega N$. La distance $\delta_{GV}(R)$, appelée distance de Gilbert-Varshamov, est définie comme le plus petit ω tel que $H(\omega) = 1 - R$. Les inégalités de Markov et de Chebychev permettent de montrer :

Si $\omega < \delta_{GV}$ ou $\omega > 1 - \delta_{GV}$, $n_w = 0$ p.s.,

Si $\omega \in [\delta_{GV}, 1 - \delta_{GV}]$, $s(\omega) \doteq \lim_{N \rightarrow \infty} \frac{1}{N} \log n_w = H(\omega) + R - 1$ p.s.

La fonction $s(\omega) = H(\omega) + R - 1$, représentée figure 1.6, s'apparente ici à une entropie microcanonique : en effet l'énergie normalisée (1.53) d'un mot de code s'écrit $e = h\omega$. L'opération de décodage, quant à elle, s'inscrit dans le cadre de l'ensemble canonique, cf. (1.52). L'équivalence des ensembles peut donc s'appliquer, une fois éclaircies certaines particularités de notre système. Notre entropie microcanonique $s(\omega)$ compte seulement les « mauvais » mots de code. Elle est donc associée à la mesure totale des erreurs, définie par :

$$Z_{\text{err}}(\beta) = \sum_{\sigma \in \mathcal{C}'} 2^{-\beta E(\sigma)} \asymp \int_{\delta_{GV}}^{1 - \delta_{GV}} d\omega 2^{N[s(\omega) - \beta h\omega]} \asymp 2^{N \max_{\omega \in [\delta_{GV}, 1 - \delta_{GV}]} (s(\omega) - \beta h\omega)}. \quad (1.64)$$

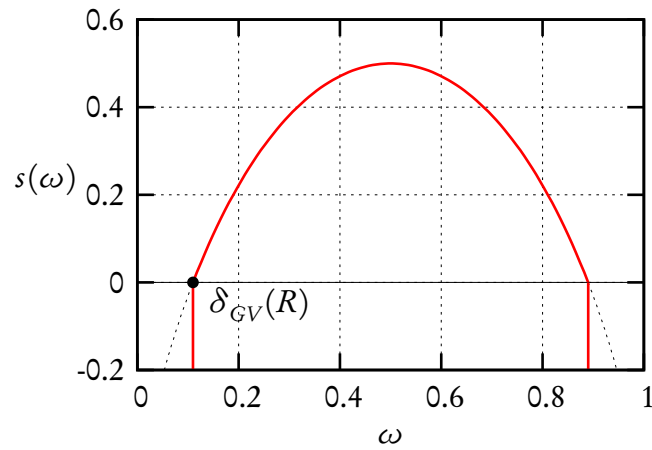


Fig. 1.6: Énumération des distances dans un code aléatoire de taux $R = 1/2$.

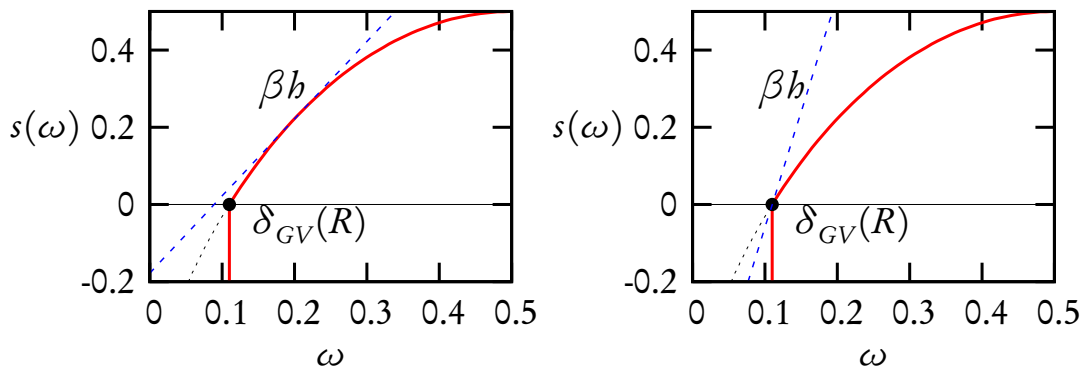


Fig. 1.7: Estimation graphique de l'erreur dans la phase « liquide » (panneau de gauche) et dans la phase condensée (panneau de droite). La transformée de Legendre se fait normalement à haute température. Au point de condensation donné par (1.67), la droite de pente βh prend appui sur le point singulier $\omega = \delta_{GV}(R)$.

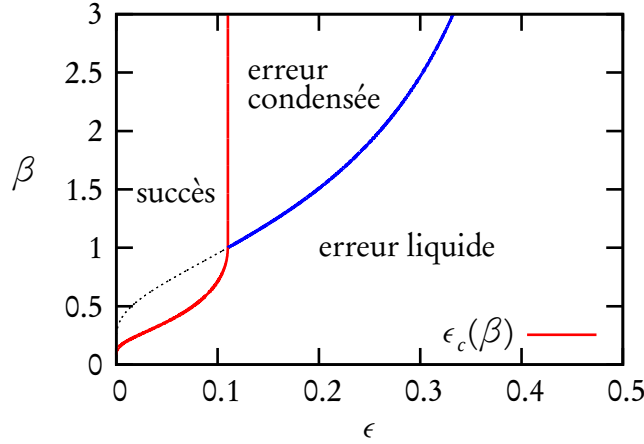


Fig. 1.8: Diagramme de phase du décodage d'un code aléatoire ($R = 1/2$) sur le BSC. Dans la phase décodable, la ligne pointillée marque la transition liquide/condensat de la phase d'erreur (dominée).

Dans l'estimation du col, il est important de remarquer que le « max » en ω peut être atteint à la frontière de l'intervalle $[\delta_{GV}, 1 - \delta_{GV}]$, où la dérivée de la fonction est non-nulle.

D'un autre côté, le poids du mot de code transmis s'exprime ainsi :

$$Z_0(\beta) = 2^{-\beta h \|\sigma_0 - \tau\|} \asymp 2^{-N\beta h \epsilon} \text{ presque sûrement.} \quad (1.65)$$

Le décodage sera un succès si et seulement si Z_0 domine exponentiellement Z_{err} . Autrement dit, si

$$g(\epsilon, \beta) \doteq \max_{\omega > \delta_{GV}(R)} [s(\omega) - \beta h \omega] + \beta h \epsilon < 0. \quad (1.66)$$

Cette inégalité peut être testée graphiquement grâce à la construction géométrique de la transformée de Legendre (cf. figure 1.7). Quand l'erreur est dominée par le point-frontière $\omega = \delta_{GV}(R)$, où l'entropie est nulle, on parle de phase *condensée*² : le poids de Boltzmann associé à l'erreur se concentre alors sur un petit nombre de mots parasites. Cela se produit si et seulement si :

$$\left(\frac{1-\epsilon}{\epsilon}\right)^\beta > \frac{1-\delta_{GV}(R)}{\delta_{GV}(R)} \quad (1.67)$$

Cette inégalité marque la séparation entre la phase condensée et une phase « liquide » dominée par un grand nombre de mots de codes. Dans la phase condensée le critère

²Dans le langage de la physique, ce phénomène par lequel l'entropie devient artificiellement négative s'appelle la *crise entropique*.

(1.66) se simplifie :

$$g(\epsilon, \beta) = \beta h \epsilon - \beta h \delta_{GV}(R). \quad (1.68)$$

Quand $\beta \geq 1$, la transition de la décodabilité $g < 0$ vers la non-décodabilité $g > 0$ se produit toujours dans la phase condensée, en $\epsilon_c(\beta) = \delta_{GV}(R)$. C'est précisément le seuil de Shannon :

$$R = C_{\text{BSC}}[\epsilon_c(\beta)] = 1 - H(\epsilon_c(\beta)). \quad (1.69)$$

Ainsi, dans toutes les situations intermédiaires entre le décodage par mot ($\beta = \infty$) et celui par lettre ($\beta = 1$), les codes aléatoires saturent encore une fois la borne de Shannon.

En revanche, si $\beta < 1$, la transition se fait dans la phase liquide. Le bruit critique $\epsilon_c(\beta) < \delta_{GV}(R)$ vérifie alors :

$$H\left(\frac{1}{1 + 2^{\beta h}}\right) + R - 1 + \beta h(\epsilon - \omega) = 0. \quad (1.70)$$

La figure 1.8 résume les différentes phases en présence dans le diagramme (ϵ, β) .

Ainsi, lorsqu'on décode à une température non-physique ($\beta < 1$), l'erreur s'explique comme la somme des effets d'un grand nombre $2^{Ns(\omega)}$ de mots de codes parasites. Au rebours, dans le cas physique $\beta \geq 1$, l'erreur vient d'un petit nombre de mauvais mots de codes. Ce phénomène de condensation, consécutif d'une discontinuité de la dérivée de l'entropie microcanonique, est caractéristique des systèmes vitreux, et trouve une illustration exemplaire dans le REM de Derrida [Der80, Der81], dont le modèle à codes aléatoires est une variante.

1.2.3 Compression avec perte

Le théorème de Shannon sur la correction d'erreur est en fait plus général que nous ne l'avons énoncé. Il indique également quel est le plus petit taux d'erreur qu'on peut espérer lorsque le taux R du code dépasse la capacité du canal. Nous en examinons ici un cas particulier, quand le canal est non-bruité ($C = 1$) : c'est le problème de la compression avec perte. On veut transformer une chaîne binaire τ de longueur M en une chaîne σ *plus courte* de longueur N . Le taux du code est redéfini par :

$$R = \frac{N}{M} < 1. \quad (1.71)$$

Par souci de compatibilité avec le formalisme des codes de compression à tests de parité, que nous discuterons au chapitre 4, les notations ont été modifiées par rapport au paragraphe précédent.

Le code est caractérisé par une fonction de codage f , qui a un mot de M lettres associe un mot de N lettres, et par une fonction de décodage g , qui effectue l'opération

inverse. Ces fonctions doivent vérifier $f[g(\sigma)] = \sigma$. L'image de la fonction g définit l'ensemble $\mathcal{C}$ des mots de code, c'est-à-dire les messages τ qui seront compressés sans perte. Pour les autres, la *distorsion* est définie par :

$$D(\tau) = \|\tau - g[f(\tau)]\|, \quad (1.72)$$

c'est-à-dire le nombre d'erreurs que subit la chaîne τ à la suite des opérations de codage et de décodage. Le théorème de Shannon affirme qu'il est possible d'atteindre le niveau de distorsion moyen suivant :

$$\lim_{M \rightarrow \infty} \frac{\mathbb{E}[D(\tau)]}{M} = \delta_{GV}(R). \quad (1.73)$$

Cette performance peut être atteinte asymptotiquement par un code aléatoire, défini comme suit : la fonction de décodage g est tirée au hasard en choisissant de manière uniforme 2^N mots de codes de longueur M parmi 2^M possibles. La fonction de codage f met en œuvre un principe d'optimisation :

$$f(\tau) = \operatorname{argmax}_{\sigma} \|g(\sigma) - \tau\|, \quad (1.74)$$

en minimisant la distorsion :

$$D(\tau) = \max_{\sigma} \|g(\sigma) - \tau\|. \quad (1.75)$$

Ce modèle aléatoire est quelque sorte le dual de celui que nous avons introduit pour la correction d'erreur. Dans la correction d'erreur la fonction de *codage* f était tirée au hasard alors que ce sort est maintenant réservé à la fonction de *décodage* g . Corrélativement, le codage constitue dans le cas présent la tâche la plus ardue, à l'inverse de la correction d'erreur, où le décodage était le plus difficile.

En vertu de (1.75), la distorsion est donnée par la distance de τ à l'ensemble des mots de code $\mathcal{C}$. Or l'analyse du spectre de distances d'un ensemble de mots aléatoires effectuée au paragraphe précédent a permis de montrer que cette distance vaut presque sûrement :

$$D(\tau) \sim M \delta_{GV}(R), \quad (1.76)$$

réalisant ainsi la borne de Shannon (1.73).

Dans tous les cas que nous avons considérés, en correction d'erreur comme en compression, les codes aléatoires font preuve de leur optimalité. Ces performances théoriques exceptionnelles sont toutefois contrebalancées par un fâcheux revers de médaille, à savoir l'impossibilité pratique de leur mise en œuvre. En effet presque toutes les tâches impliquées dans leur fonctionnement requièrent un nombre exponentiellement élevé ($\approx 2^N$) d'opérations, rendant illusoire l'espoir d'une utilisation réelle. Heureusement, nous verrons au chapitre 4 que des constructions de codes linéaires et dilués offrent une alternative praticable aux codes aléatoires.

1.2.4 Effets de taille finie

Les résultats discutés jusqu'ici concernent les performances de codes dans la limite des longs mots. Mais en pratique les codes sont évidemment limités en longueur. Comment estimer les effets de taille finie, et quelles conséquences en tirer en termes de fiabilité ? Dès les années 50 et 60, peu après l'émergence des fondations de la théorie de l'information, Shannon et certains de ses collègues tentèrent de répondre à cette question. Notamment, de nombreux efforts furent consacrés à la recherche de la meilleure performance possible en longueur finie, qui correspondrait à une sorte de borne de Shannon de la performance asymptotique. Il s'agit dans ce problème de minimiser la probabilité d'erreur $P_e(\mathcal{C})$ sur les codes $\mathcal{C}$ de longueur N et de taux R . Il a pu être montré que pour les meilleurs codes cette probabilité décroît exponentiellement avec la taille :

$$P_e(\mathcal{C}_N) \asymp 2^{-NE(R)}, \quad (1.77)$$

où $\{\mathcal{C}_N\}_{N>0}$ est une séquence de codes de longueur croissante. Pour une séquence donnée, l'exposant E est appelé *exposant d'erreur*. Le plus grand exposant possible définit la *fonction de fiabilité* du canal. Le calcul de cette fonction est en fait beaucoup plus difficile que celui de la borne de Shannon : bien qu'elle ait pu être caractérisée dans certains cas, et que de nombreuses bornes aient été proposées, son évaluation générale est encore l'objet de conjectures.

Nous ne traitons pas ici directement la question de la fonction de fiabilité, et concentrons notre analyse sur les exposants d'erreur de séquences aléatoires. Par souci de simplicité, les probabilités seront ici considérées à la fois par rapport aux réalisations du bruit et au choix du code, que l'on supposera tiré d'un ensemble particulier. Le mélange de ces deux niveaux de désordre conduit à l'évaluation d'exposants d'erreur *moyens*, par opposition aux exposants d'erreurs *typiques*. Formellement, ces deux types d'exposant sont définis par :

$$E_{\text{moy}} = \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E}_{\mathcal{C}} P_e(\mathcal{C}_N), \quad (1.78)$$

$$E_{\text{typ}} = \lim_{N \rightarrow \infty} \frac{1}{N} \mathbb{E}_{\mathcal{C}} \log P_e(\mathcal{C}_N), \quad (1.79)$$

où $\mathbb{E}_{\mathcal{C}}$ désigne la moyenne sur un ensemble donné de codes de longueur N .

Nous utilisons un formalisme de grandes déviations basés sur les quantités thermodynamiques que nous avons introduites dans le cadre des codes aléatoires : par exemple, dans le canal d'effacement, n désigne le nombre de mots de codes compatibles avec le mot reçu, autres que le mot original. On écrit $n = 2^{N^s}$, où s est interprétée comme une entropie. Nous postulons que cette fonction d'entropie est, pour une séquence de constructions de codes et de réalisations du bruit, soumises à un principe

de grandes déviations :

$$\mathbb{P}\left(\frac{1}{N} \log n = s\right) \asymp 2^{NL(s)} \quad (1.80)$$

La *fonction de taux* $L(\cdot)$ est toujours négative ou nulle, et elle est maximale en la valeur typique de ses arguments, où elle vaut 0. Si nous nous plaçons par exemple dans la phase décodable du canal d'effacement, $\epsilon < 1 - R$, alors L est maximale en $s = -\infty$. Le cas où le décodage échoue, $s \geq 0$, correspond alors à des réalisations atypiques, donc exponentiellement improbables, du hasard. L'exposant d'erreur s'exprime ainsi :

$$\text{BEC} \quad E_{\text{moy}} = -\sup_{s \geq 0} L(s). \quad (1.81)$$

Dans la phase décodable ($\epsilon < 1 - R$), la loi de probabilité du nombre de mauvais mots de code s'écrit :

$$\mathbb{P}(n) = \sum_{|E|=0}^N \binom{N}{|E|} \epsilon^{|E|} (1-\epsilon)^{N-|E|} \mathbb{P}(n \text{ sachant } |E|) \quad (1.82)$$

où $|E|$ est le nombre de bits effacés, et où $\mathbb{P}(n \text{ sachant } |E|)$ suit une loi binomiale :

$$\mathbb{P}(n \text{ sachant } |E|) = \binom{2^L - 1}{n} 2^{n(|E|-N)} (1 - 2^{|E|-N})^{2^L - 1 - n} \quad (1.83)$$

(chaque mot de code est, indépendamment des autres, compatible avec le bruit E avec probabilité $2^{|E|-N}$: il doit coïncider avec le mot reçu sur le domaine $\{1, \dots, N\} \setminus E$ des bits correctement transmis).

Le comportement asymptotique de cette probabilité conditionnelle diffère suivant l'intensité du bruit. On note $|E| = eN$, et $n = 2^{Ns}$. Si $e > 1 - R$, alors l'échec ($s \geq 0$) est presque sûr. Si au contraire $e < 1 - R$, alors :

$$\frac{1}{N} \log \mathbb{P}(n = 2^{Ns} | |E| = Ne) \sim (R - 1 + e - s) 2^{Ns}, \quad (1.84)$$

d'où $\mathbb{P}(n > 0 | |E| = Ne) \asymp 2^{N(R-1+e)}$.

En résumant :

$$\mathbb{P}(n > 0) \asymp \int_0^{1-R} de \, 2^{-N[D(e|\epsilon)+1-R-e]} + \int_{1-R}^1 de \, 2^{-ND(e|\epsilon)}, \quad (1.85)$$

soit

$$E_{\text{moy}} = \inf_{e \in [0,1]} [D(e|\epsilon) + \max(0, 1 - R - e)]. \quad (1.86)$$

Deux régimes apparaissent suivant le niveau du bruit ϵ et le taux R . Près de la borne de Shannon, pour $\epsilon > (1 - R)/(1 + R)$, l'infimum est atteint en $e = 1 - R$. Cela signifie

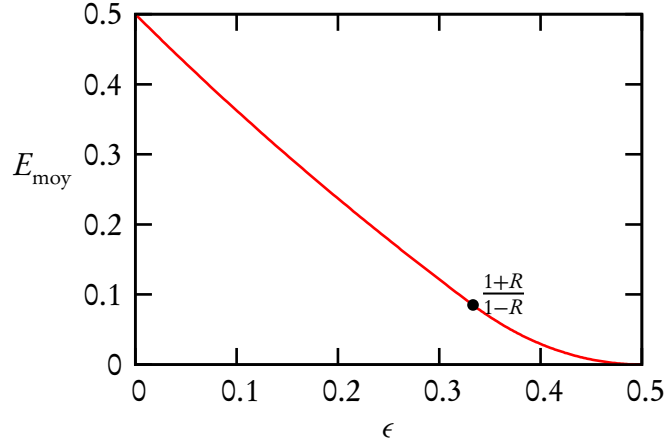


Fig. 1.9: Exposant d'erreur en fonction de la probabilité d'effacement, pour $R = 1/2$.

que l'erreur est typiquement causée par une réalisation du bruit exceptionnellement défavorable, au point d'atteindre la borne de Shannon. Plus loin de cette borne, pour $\epsilon < (1-R)/(1+R)$, l'infimum est atteint en $e^* < 1-R$. Là, l'erreur est dominée par la conjonction de deux événements improbables : un bruit important et l'existence exceptionnelle d'un mot de code compatible à ce niveau de bruit. On obtient (voir figure 1.9) :

$$E_{\text{moy}} = \begin{cases} D(1-R||\epsilon) & \text{si } \epsilon \in [(1-R)/(1+R), 1-R] \\ 1-R - \log(1+\epsilon) & \text{sinon} \end{cases} \quad (1.87)$$

Incidentement, le taux de grande déviation vaut, pour $s \geq 0$:

$$L(s) = \begin{cases} -D(1-R+s||\epsilon) & \text{si } \epsilon \in [(1-R)/(1+R), 1-R] \text{ ou } s > 0 \\ R-1 + \log(1+\epsilon) & \text{sinon} \end{cases} \quad (1.88)$$

et $L(s = -\infty) = 0$. La figure 1.10 représente la fonction de taux dans les deux situations décrites plus haut.

Le cas du canal symétrique BSC se traite avec des arguments similaires, en définissant une fonction de grande déviation portant sur les énergie libres d'erreur et de succès, associées à Z_{err} et à Z_0 . On y observe également deux régimes suivant le niveau moyen de bruit ϵ et le taux R , avec des interprétations homologues : près de la borne de Shannon, l'erreur est principalement causée par un bruit anormalement élevé ; lorsqu'on s'éloigne de cette borne, l'apparition d'un mot de code anormalement proche du mot reçu rentre aussi en jeu.

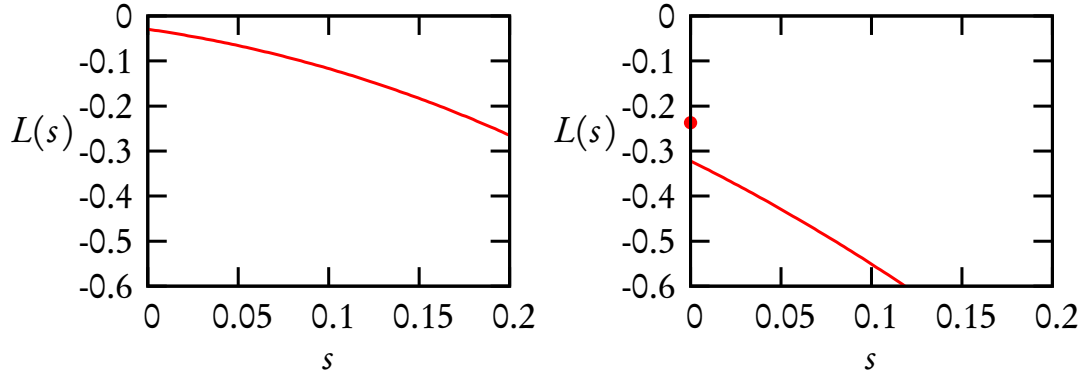


Fig. 1.10: Taux de grande déviation $L(s)$ pour $\epsilon = 0,4 > (1 - R)/(1 + R)$ (panneau de gauche) et $\epsilon = 0,2 < (1 - R)/(1 + R)$ (panneau de droite), avec $R = 1/2$. Dans le deuxième cas, le taux est discontinu en $s = 0$.

Références

L'article fondamental de Shannon [Sha48] constitue une bonne introduction à la théorie de l'information. Cet article s'avère en fait particulièrement pédagogique et facile d'accès, et fournit de nombreux exemples et illustrations. Un livre classique traitant de la théorie de l'information est celui de Cover et Thomas [CT91]. Le livre plus récent de MacKay [Mac03] fournit un exposé très clair de nombreux sujets en théorie de l'information, et fait le lien avec les domaines de l'inférence et de l'apprentissage. Pour une bonne introduction à l'entropie physique et ses relations avec l'entropie de Shannon on peut consulter [Bal83] ou [DGLR89]. L'identification des problèmes de codes de correction d'erreur à des modèles de verres de spin a été pointée par Sourlas [Sou89, Sou94]. La formulation physique du modèle à codes aléatoires est due à Montanari [Mon01] et est reprise dans un livre en préparation [MM07]. Le problème de la fonction de fiabilité est discuté dans [Ber02], et certains calculs d'exposants d'erreur sont présentés dans [BJ02]. Les appendices de [MR06a] reprennent ces calculs sur les canaux symétriques et d'effacement, et introduisent le formalisme thermodynamique des exposants d'erreur.

Les éléments de théorie de l'information exposés ici introduisent les concepts généraux utiles à la lecture de [MR06a, MR06b]. Ils sont également l'occasion de présenter certains outils importants de la physique statistique, en particulier les notions d'ensemble thermodynamique, de transformée de Legendre et de transition vitreuse.

Chapitre 2

Approche physique de la complexité

Ce chapitre aborde les problèmes d'optimisation et de satisfaction de contraintes en adoptant le point de vue de la physique statistique. Après un bref passage en revue de la théorie classique de la complexité, les notions et apports importants de la physique statistique sont exposés et illustrés sur le problème de la satisfaisabilité.

2.1 Théorie classique de la complexité

2.1.1 Optimisation combinatoire

Un problème d'optimisation combinatoire est défini par une application qui à chaque configuration d'un espace discret associe un *coût* :

$$\mathcal{C} \longrightarrow \mathbb{R} \quad (2.1)$$

$$\sigma \longmapsto E(\sigma) \quad (2.2)$$

Il s'agit de trouver la configuration de plus bas coût, c'est-à-dire σ^* tel que :

$$E(\sigma^*) \leq E(\sigma) \quad \forall \sigma \in \mathcal{C}. \quad (2.3)$$

La procédure de décodage par mots sur le canal binaire symétrique, décrite au chapitre précédent (paragraphe 1.2.1), fournit un exemple de tel problème. L'espace configurationnel est alors défini comme l'ensemble des mots de codes $\mathcal{C}$, et la fonction à minimiser est la distance de Hamming au mot reçu.

Le problème du voyageur de commerce est l'un des exemples les plus cités (et les plus simples à formuler) de problème d'optimisation. Étant données une liste de villes indicées par $i \in \{1, \dots, N\}$ et un jeu de distances entre elles $\{d_{ij}\}$, quel est le parcours le moins long qui passe par chacune des villes exactement une fois? Un parcours est

formalisé par une permutation à N éléments, $\sigma \in \mathcal{S}_N = \mathcal{C}$, et la fonction de coût s'écrit comme la somme des longueurs de chaque étape :

$$\sum_{i=1}^{N-1} d_{\sigma_i, \sigma_{i+1}}. \quad (2.4)$$

Satisfaction de contraintes

Les problèmes de satisfaction de contraintes peuvent être vus comme une sous-classe des problèmes d'optimisation. Sur un espace de configurations $\mathcal{C}$ on impose un certain nombre de contraintes logiques c . La collection de ces contraintes forme une *instance*, notée $\mathcal{F}$. Il s'agit alors de trouver une configuration qui satisfasse toutes les contraintes de l'instance. À chaque instance de contraintes on peut associer une fonction de coût, qui compte le nombre de contraintes violées :

$$E(\sigma) = \sum_{c \in \mathcal{F}} \mathbb{I}(\sigma \not\models c) \quad (2.5)$$

où « $\sigma \models c$ » signifie que la configuration σ satisfait la contrainte c . Résoudre le problème de satisfaction de contraintes revient alors à trouver une configuration de coût nul. Savoir si une telle configuration existe relève d'un problème de *décision*. Cette question est en principe plus aisée que celle de l'optimisation. Cependant, pour chaque problème d'optimisation combinatoire on peut se poser la question (décisionnelle) de l'existence ou non d'une configuration de coût inférieur à une valeur donnée. Par une méthode de dichotomie, on peut ainsi ramener un problème d'optimisation à une succession de problèmes de décision bien choisis.

L'exemple le plus célèbre et le plus étudié de problème de satisfaction de contraintes est celui de la *satisfaisabilité*, souvent abrégé en SAT. Les configurations y sont des chaînes de variables booléennes $\sigma_i \in \{\text{VRAI}, \text{FAUX}\}$, $i = 1, \dots, N$, aussi appelées valuations de vérité. Chaque contrainte, ou *clause*, est une fonction disjonctive de littéraux (variables booléennes ou leur négation). Par exemple :

$$\neg\sigma_2 \text{ OU } \sigma_5 \text{ OU } \neg\sigma_7. \quad (2.6)$$

Mises bout à bout et reliées par des ET, ces clauses constituent une formule logique sous *forme normale conjonctive*. S'il existe une valuation des variables qui rend cette formule vraie, elle est dite *satisfaisable* (SAT) ; sinon, elle est insatisfaisable (non-SAT). Le problème SAT tire sa généralité du fait que toute formule de logique booléenne peut effectivement s'écrire sous forme normale conjonctive. La version « optimisation » du problème de satisfaisabilité, qui consiste à minimiser le nombre de clauses violées, est communément appelée « maximum SAT », ou « MAX-SAT ».

Le problème SAT a surtout été étudié dans le cadre de la logique booléenne, où il joue un rôle majeur dans les problèmes de vérification, de contrôle de modèle (*model*

checking), ou d'automatisation de preuves. Plus généralement, la satisfaisabilité intervient dans de nombreux domaines de l'informatique, que ce soit en algorithmique, en intelligence artificielle (planification, diagnostic) ou en conception industrielle (traitement pipeline des processeurs).

2.1.2 P vs NP

Le problème de satisfaisabilité est également une des pierres angulaires de la théorie de la complexité algorithmique, que nous esquissons ici sommairement. Étant donnée une question décisionnelle, telle que celle de la satisfaisabilité d'une formule donnée, combien d'opérations faut-il à un ordinateur pour y répondre ? Ici, un « ordinateur » est formellement une machine de Turing, c'est à dire une machine capable d'effectuer séquentiellement des opérations logiques sur des variables discrètes stockées en mémoire. S'il est vrai que tous les problèmes de logique booléenne sont solubles par un ordinateur en un temps — i.e. nombre d'opérations — fini, l'ordre de grandeur de ce temps peut s'avérer prohibitif pour les applications pratiques. Par exemple, si nous voulons déterminer si une formule logique est SAT, nous pouvons tester chacune des 2^N valuations possibles des variables booléennes et constater si elle satisfait ou non la formule. Mais pour les longues formules, le nombre d'opérations requise par cet algorithme devient très vite énorme. La question de l'existence d'algorithmes plus performants apparaît donc ici d'une importance cruciale.

Afin de formaliser la notion de complexité algorithmique, introduisons quelques définitions. La taille d'une instance d'un problème discret désigne le nombre de symboles nécessaires pour décrire cette instance. Un problème est dit *polynomial* (abbrev. P) si *chacune* des instances de ce problème peut être résolue par un algorithme polynomial, i.e. dont le nombre d'opérations est borné uniformément par un polynôme prenant comme argument la taille de l'instance. La classe NP , ou *non-déterministe polynomiale*, contient plus largement les problèmes décisionnels solubles par un algorithme polynomial exécuté par une machine de Turing non-déterministe. En clair, cela veut dire qu'une configuration σ candidate à la question de décision peut être testée par un algorithme polynomial. Beaucoup des problèmes que nous rencontrons dans cette thèse appartiennent à cette dernière catégorie. Malheureusement, il est aujourd'hui communément admis que tous les problèmes NP ne peuvent pas être résolus par un algorithme polynomial : c'est la célèbre conjecture $P \neq NP$, qui reste indémontrée à ce jour. Néanmoins, un résultat important dû à Cook [Coo71] affirme que toute instance de problème appartenant à la classe NP peut être traduite en une instance SAT par un algorithme polynomial. SAT est ainsi *au moins aussi dur* que n'importe quel autre problème dans NP : on le dit « NP -complet ». Depuis ce résultat précurseur, un grand nombre de problèmes ont été identifiés comme étant NP -complets. Les problèmes NP -complets tirent leur importance du fait que la preuve

de leur appartenance à P entraînerait immédiatement $P = NP$. Dans l'état actuel des connaissances, les meilleurs algorithmes traitant exhaustivement des problèmes NP -complets sont exponentiels dans la taille du problème, laissant peu d'espoir de résoudre systématiquement les instances de grande taille.

Les classes de complexités que nous venons de décrire peuvent s'illustrer par des cas spéciaux de SAT. On définit le problème k -SAT comme l'ensemble des formules de satisfaisabilité dont chaque clause contient exactement k littéraux. On peut démontrer que 2-SAT est polynomial, alors que k -SAT est NP -complet pour chaque $k \geq 3$ (cf. [Pap94], §9.2). Cette observation suggère l'existence d'une frontière entre les problèmes « faciles » (polynomiaux) et les problèmes « difficiles » (NP -complets, supposément exponentiels), et explicite cette frontière dans le cas particulier de k -SAT. Cependant, la pertinence d'une telle dichotomie reste problématique, comme en témoigne l'exemple du problème d'isomorphisme de graphes, pour lequel aucun algorithme polynomial ni aucune preuve de NP -complétude n'ont pu être trouvés. Plus rigoureusement, Ladner [Lad75] a démontré qu'il existe, sous l'hypothèse $P \neq NP$, des problèmes qui ne sont ni NP -complets ni polynomiaux.

2.2 Complexité du pire et complexité typique

2.2.1 Motivation

Notre discussion sur la complexité algorithmique peint jusqu'ici un tableau plutôt sombre des capacités de résolution algorithmique de problèmes complexes tels que celui de la satisfaisabilité. En pratique cependant, il est souvent possible de résoudre en un temps raisonnable de larges classes d'instances de problèmes NP -complets, à l'aide d'algorithmes simples. Parmi les plus célèbres des algorithmes complets de résolution de SAT, on compte celui de Davis-Putman-Logemann-Loveland (DPLL) [DLL62], ainsi que tous ses dérivés. Il existe également un large choix d'algorithmes stochastiques, tels que *Random WalkSAT* [Pap91], capables de trouver des solutions rapidement, mais impropres à certifier la non-satisfaisabilité. Pour beaucoup d'instances, ces algorithmes donnent une réponse en un nombre raisonnable d'opérations.

Rappelons que la théorie classique de la complexité repose sur la performance d'algorithmes *uniformément sur toutes les instances*. La plupart du temps, cette performance est en fait limitée par un petit nombre de « mauvaises » instances. L'acceptation du qualificatif « mauvais » dépend elle-même de l'algorithme considéré, ce qui rend ardue la définition précise d'une difficulté intrinsèque.

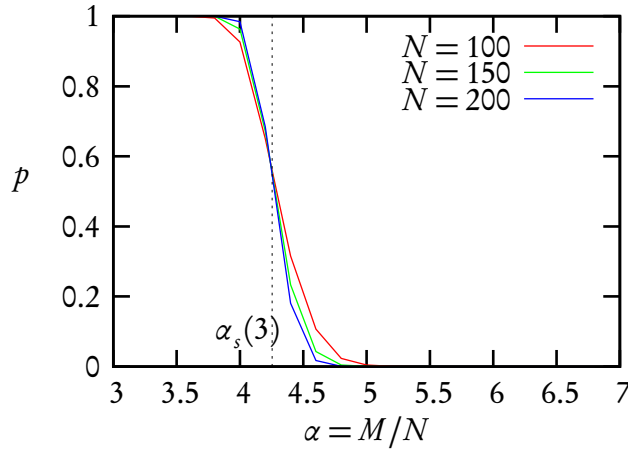


Fig. 2.1: Probabilité de satisfaisabilité pour le problème 3-SAT en fonction de α , pour diverses valeurs de N . Dans la limite des grandes formules, cette probabilité connaît un seuil abrupt en $\alpha_s(3) \approx 4,26$.

2.2.2 Ensembles aléatoires et transitions de phase

Ainsi, la théorie classique de la complexité s'intéresse à la complexité « du pire des cas ». Cette limitation conduit naturellement à s'interroger sur la complexité des instances « typiques ». Mais qu'entend-on exactement par typique ? Afin de donner un sens à cette notion, nous devons considérer des *ensembles* d'instances aléatoires, définis par une mesure de probabilité sur les formules possibles. Un des exemples les plus simples d'ensemble, qui retiendra particulièrement notre attention, est l'ensemble k -SAT aléatoire, où N variables booléennes sont soumises à M clauses de k littéraux, chacune étant tirée uniformément parmi les $2^k \binom{N}{k}$ possibles. Dans la limite où N et M tendent vers l'infini tout en maintenant constante la densité de clauses $\alpha = M/N$, les formules « typiques » forment une sous-partie majoritaire d'instances (formellement, une séquence de sous-parties dont la mesure totale tend vers 1). Il est toutefois important de remarquer qu'en pratique la typicité ne prend de sens qu'en regard d'une propriété précise, car il existe une infinité de façons de prendre une sous-partie majoritaire.

Examinons la propriété de satisfaisabilité pour l'ensemble k -SAT aléatoire : dans la limite des grandes instances, quelle est la probabilité qu'une formule aléatoire soit satisfaisable ? Il va de soi que cette probabilité doit décroître avec α , car l'ajout de nouvelles clauses ne peut que diminuer les chances de succès. En fait, ainsi que l'illustre la figure 2.1, cette probabilité tend vers une marche d'escalier quand $N \rightarrow \infty$. Quand la densité $\alpha = M/N$ est inférieure à un seuil critique $\alpha_s(k)$, les formules aléatoires

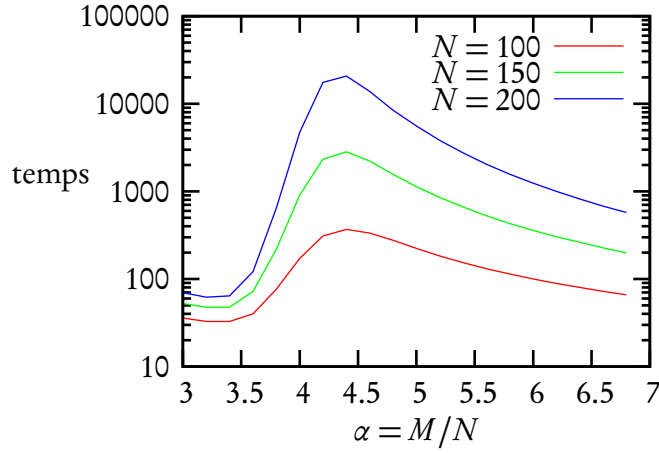


Fig. 2.2: Taille moyenne de l'arbre exploré par l'algorithme DPLL en fonction de α pour diverses valeurs de N . Cette taille est proportionnellement reliée au nombre d'opérations requises pour trouver une solution ou prouver l'insatisfaisabilité.

sont presque sûrement satisfaisables¹ ; à l'inverse, quand $\alpha > \alpha_s(k)$, elle sont presque sûrement insatisfaisables : la propriété de satisfaisabilité subit une transition *abrupte*.

Friedgut a pu montrer l'existence d'une transition abrupte dans k -SAT [Fri99], et son argument a pu être étendu à d'autres problèmes présentant des caractéristiques similaires. Son résultat implique l'existence d'un seuil non uniforme $\alpha_N(k)$, tel que :

$$\forall \epsilon, \quad \lim_{N \rightarrow \infty} \mathbb{P}(\text{satisfaisable}) = \begin{cases} 1 & \text{si } \alpha < \alpha_N(k)(1 - \epsilon), \\ 0 & \text{si } \alpha > \alpha_N(k)(1 + \epsilon), \end{cases} \quad (2.7)$$

La convergence de $\alpha_N(k)$ n'est cependant pas garantie, et reste à ce jour à l'état de conjecture. Des bornes rigoureuses ont toutefois pu être établies en utilisant des méthodes de premier et de second moment pour les bornes supérieure [DB97, KKKS98] et inférieure [AP04], respectivement. Nous reviendrons sur les méthodes employées pour dériver ces bornes dans le chapitre 5.

Quelles performances les algorithmes classiques de SAT affichent-ils sur les problèmes aléatoires ? La figure 2.2 représente en échelle logarithmique le nombre moyen d'opérations exécutées par l'algorithme DPLL sur des instances du problème 3-SAT aléatoire. On constate que ce nombre connaît un pic autour de la transition $\alpha_s(k)$, et semble croître exponentiellement avec la taille du problème. Cette dernière observation fait l'objet d'un résultat rigoureux [CS88] dans toute la phase non-SAT. Dans la phase SAT, [CF86, CF90] démontrent que les algorithmes de type DPLL trouvent

¹On dit qu'un événement se produit *presque sûrement* quand sa probabilité tend vers 1 alors que $N \rightarrow \infty$.

une solution en temps polynomial pour des densités α suffisamment faibles. À l’opposé, quand la densité est proche du seuil, ces algorithmes requièrent un nombre exponentiel d’opérations [ABM01]. Cette seconde transition n’est pas sans rappeler le phénomène de transition vitreuse en physique, par lequel la dynamique devient soudainement très lente, et peine à trouver les minima d’énergies. De la même manière, l’algorithme DPLL reste longtemps piégé dans des régions défavorables de l’espace configurationnel avant de trouver une solution. Cette analogie contribue à motiver l’intérêt d’une approche physique des problèmes aléatoires de satisfaction de contrainte. Parallèlement, ces résultats nourrissent l’espoir de construire des formules *vraiment* difficiles autour de la transition et ainsi comprendre la nature distinctive des problèmes NP -complets. Toutefois, ces résultats concernent DPLL et ne préjugent en rien de la capacité à résoudre efficacement les formules aléatoires difficiles à l’aide d’algorithmes incomplets, comme en témoignent les performances des algorithmes *Random Walk-SAT* et *Survey Inspired Decimation* [MPZ02].

2.2.3 Les problèmes réels sont-ils aléatoires ?

Le choix de l’ensemble k -SAT aléatoire comme cadre de référence d’une théorie de la complexité typique soulève un certain nombre d’objections. En particulier, les problèmes aléatoires sont assez éloignés des problèmes réels, dans lesquels les effets potentiels de concentration sont plus fréquents. La concentration se caractérise par l’existence d’un certain nombre de contraintes dont les desiderata (c’est-à-dire les littéraux) sont partiellement contradictoires. Alors que ces effets sont fréquents dans les problèmes réels, ils sont presque totalement absents des ensembles aléatoires que nous avons décrits. Considérons par exemple la probabilité que deux variables soient toutes deux impliquées dans plusieurs clauses à la fois. Dans 3-SAT aléatoire, cette probabilité se comporte comme $18\alpha^2/N^2$ pour chaque paire de variables ; la proportion de telles paires tend donc vers zéro quand N tend vers l’infini, contrairement à ce qu’on observe dans les problèmes réels. Cet exemple appuie le constat selon lequel les problèmes aléatoires constituent une classe bien particulière d’instances, qui pourraient en fait s’avérer « anormalement » faciles. En effet, ainsi que nous l’avons déjà souligné, il n’est pas exclu que certains problèmes NP -complets admettent des sous-classes contenant la plupart des instances et solubles par des algorithmes polynomiaux². La question de savoir si k -SAT aléatoire produit une telle sous-classe reste toutefois ouverte, ainsi que l’est la possibilité de construire des instances véritablement difficiles à partir de cet ensemble.

²C’est le cas par exemple du problème 1-parmi- k -SAT [ACIM01].

2.3 Diagramme de phases

2.3.1 Formulation physique

L'approche physique des problèmes de satisfaction de contrainte et d'optimisation repose sur la définition d'une mesure de Boltzmann :

$$p_\beta(\sigma) = \frac{2^{-\beta E(\sigma)}}{Z(\beta)}, \quad (2.8)$$

où $E(\sigma)$ est la fonction de coût du problème d'optimisation — ou encore le nombre de contraintes violées du problème de satisfaction de contraintes. La solution du problème est donnée par la limite de température nulle ($\beta \rightarrow \infty$), qui permet d'accéder au fondamental. Dans le cas des problèmes de satisfaction de contrainte *satisfaisables*, la limite d'énergie nulle s'écrit :

$$p(\sigma) = \frac{1}{Z} \prod_{a \in \mathcal{F}} \mathbb{I}(\sigma \models a). \quad (2.9)$$

où a est l'une des contraintes de l'instance $\mathcal{F}$. Ainsi, la mesure est uniformément répartie sur l'ensemble des solutions du problèmes, et la fonction de partition Z compte le nombre total de solutions.

Le calcul de l'espérance de $Z(\beta)$ dans un ensemble donné d'instances (par exemple, k -SAT aléatoire) ne fournit que peu d'information sur sa distribution. En effet, $Z(\beta)$ se comporte comme une variable aléatoire multiplicative (elle croît exponentiellement avec N), et ne jouit pas d'une propriété d'automoyennage. En revanche, l'énergie libre $F(\beta) = -\beta^{-1} \log Z(\beta)$ ou l'entropie du fondamental $S = \log Z$ vérifient souvent cette propriété, à savoir :

$$\mathbb{P} \left[\left| \frac{F(\beta) - \mathbb{E}F(\beta)}{N} \right| > \epsilon \right] \rightarrow 0 \quad \forall \epsilon > 0 \quad (2.10)$$

Par analogie avec la physique des systèmes vitreux, le calcul de l'espérance de $Z(\beta)$ s'appelle une moyenne *recuite*, et celui de l'espérance de l'énergie libre, une moyenne *gelée*. Dans le premier cas, les configurations et le choix de l'instance sont traités sur un pied d'égalité, car les sommes sur ces deux types de variables sont interchangeables : σ et $\mathcal{F}$ jouent le rôle de variables dynamiques. À l'opposé, en vertu de la propriété d'automoyennage, la moyenne gelée donne la fonction de partition d'une instance *typique* fixée. L'instance aléatoire $\mathcal{F}$ est identifiée à du *désordre gelé* dans le langage de la physique statistique, et seule σ a le statut de variable dynamique.

Le terme de recuit fait référence à une technique de métallurgie par laquelle un matériau est successivement réchauffé puis lentement refroidi afin de trouver le cristal

et de réduire les défauts. Le réchauffage permet de dépiéger les configurations locales défavorables en « redynamisant » certains degrés de liberté. À l’opposé, le procédé de la trempe consiste à refroidir brutalement un matériau, gelant certains degrés de liberté de manière aléatoire.

Les moyennes recuites et gelées vérifient l’égalité de convexité :

$$\mathbb{E}[\log Z(\beta)] \leq \log \mathbb{E}[Z(\beta)]. \quad (2.11)$$

Dans le but de déterminer le comportement typique d’une instance aléatoire, l’ensemble gelé sera naturellement privilégié. Malheureusement, alors que la moyenne recuite se laisse volontier calculer à l’aide de méthodes combinatoires simples, la moyenne gelée présente des difficultés liées à la présence du logarithme. Ces difficultés peuvent être surmontées au moyen de l’astuce des répliques, qui a largement fait ses preuves dans le contexte des verres de spins et problèmes assimilés [MPV87] (voir [MZ96, MZ97, BMW00] pour son application à k -SAT). Nous ne nous étendrons pas ici sur cette technique, et lui préférons la méthode de la cavité [MP01, MZ02], qui lui est équivalente et qui repose sur des hypothèses plus intuitives, tout en étant plus encline à un traitement rigoureux. Les principes de la méthode de la cavité dans le contexte des modèles graphiques (qui englobent k -SAT) seront exposés aux chapitres 4 et 6.

La moyenne recuite du nombre de solutions Z d’une instance aléatoire de k -SAT s’exprime ainsi :

$$\mathbb{E}(Z) = \mathbb{E} \left[\sum_{\sigma} \prod_a \mathbb{I}(\sigma \models a) \right]. \quad (2.12)$$

La moyenne est prise par rapport à un choix aléatoire des M clauses. Ces choix étant indépendants, les moyennes sur les fonctions indicatrices se découpent :

$$\mathbb{E}(Z) = \sum_{\sigma} \mathbb{P}(\sigma \models a)^M \quad (2.13)$$

La probabilité $\mathbb{P}_c(\sigma \models a)$ ne dépend pas de la configuration σ , et elle vaut $1 - 2^{-k}$ (un seul choix des négations parmi 2^k rend la configuration σ non-satisfaisante). On obtient donc :

$$\mathbb{E}(Z) = 2^N (1 - 2^{-k})^M \asymp 2^N [1 + \alpha \log(1 - 2^{-k})] \quad (2.14)$$

Par l’inégalité de Markov, $\mathbb{P}(Z \geq 1) \leq \mathbb{E}(Z)$, on obtient une borne supérieure [FP83] sur le seuil α_s :

$$\alpha_s \leq -\frac{1}{\log(1 - 2^{-k})} < 2^k \ln 2. \quad (2.15)$$

Cette borne supérieure s’avère en fait assez précise : il a été prouvé [AP04] qu’elle donne le bon comportement asymptotique à k grand : $\alpha_s \sim 2^k \ln(2)$. Des bornes plus

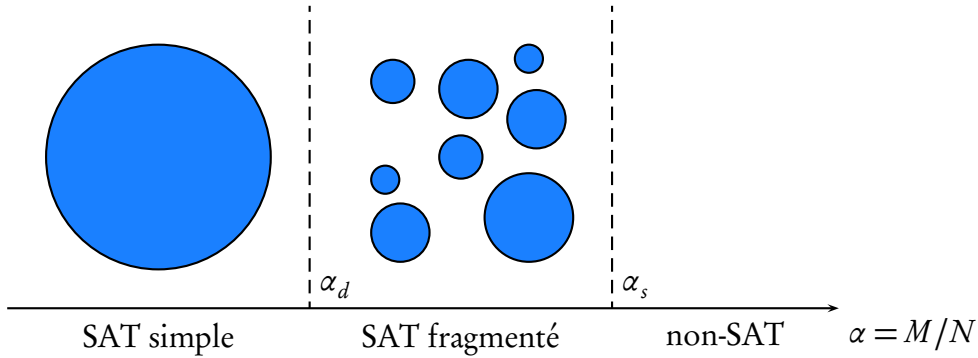


Fig. 2.3: Diagramme de phase du problème k -SAT aléatoire dans la limite $N \rightarrow \infty$. Alors que la densité de clauses $\alpha = M/N$ augmente et atteint α_d , l'espace des solutions se fragmente en un grand nombre d'amas fragmentés. Ce nombre ne cesse de décroître avant de disparaître tout-à-fait en α_s .

précises, fondées sur la même méthode du « premier moment », ont pu être dérivées [DB97, KKKS98] en considérant un ensemble plus restreint de solutions : le nombre Z n'est plus défini comme le nombre total de solution, mais comme le nombre de solutions « négativement premières ». Ces solutions se distinguent des autres par le fait qu'aucune des variables assignées à VRAI ne peut être changée en FAUX sans violer la formule. Il est facile de voir que toute formule satisfaisable admet des solutions négativement premières : il suffit pour cela de prendre n'importe quelle solution, et de changer les variables VRAI en FAUX, tant que cette opération ne conduit pas à violer la formule. L'estimation en espérance de ce nouveau Z , plus petit, réduit l'influence des formules ayant beaucoup de solutions, et livre une borne supérieure plus précise.

2.3.2 Fragmentation et condensation

Nous anticipons sur les prédictions de la méthode de la cavité en dressant un rapide tableau des propriétés intéressantes qu'elle dévoile sur l'exemple du problème k -SAT aléatoire. La plus frappante de ces propriétés est sans doute celle de la fragmentation en amas, illustrée par la figure 2.3, et intuitivement énoncée comme suit. Alors que dans le régime des basses densités de clauses, l'espace des solutions forme une grande partie connexe, cet espace se sépare, pour des densités plus élevées (mais inférieures au seuil SAT/non-SAT), en un nombre exponentiel de sous-parties connexes éloignées les unes des autres, appelées « amas ». Bien que cette caractérisation semble requérir une définition précise de la connexité, celle-ci n'a que peu d'influence sur la validité du phénomène. Pour k -SAT, il est raisonnable de définir comme *adjacentes* deux solutions ne différant que par une variable, et d'en laisser découler la notion de connexité. Pour d'autres problèmes en revanche, une acception plus souple de l'ad-

jacence devra être retenue. En général, deux solutions seront dites adjacentes si elles diffèrent par au plus $\epsilon(N)$ variables, où $\epsilon(N)$ est une fonction prescrite à l'avance et vérifiant :

$$1 \leq \epsilon(N) \leq o(N) \quad (2.16)$$

Voyons comment cette séparation en amas se manifeste sur la description de la mesure (2.9) dans la phase SAT. Appelons Z_c le nombre de solutions contenue dans un amas c , et $S_c = \log Z_c$ son entropie interne. La *complexité* ou *entropie configurationnelle* $\Sigma(s)$ mesure le nombre d'amas d'entropie donnée, supposé exponentiel :

$$\Sigma(s) = \frac{1}{N} \log \sum_c \mathbb{I}(S_c = Ns), \quad (2.17)$$

Le support de la fonction $\Sigma(s)$ est un intervalle compact dénoté $[s_m, s_M]$, au bord duquel elle s'annule.

La fonction de partition à température nulle peut alors s'écrire :

$$Z = \sum_c Z_c = \sum_c 2^{S_c} = \int_{s_m}^{s_M} ds \, 2^{N[\Sigma(s)+s]} \quad (2.18)$$

Cette expression rappelle la décomposition utilisée pour l'équivalence des ensembles, telle que nous l'avons étudiée dans le chapitre précédent, cf. (1.15). Chaque amas joue ici le rôle d'une configuration, et son entropie interne, celui d'une énergie. Par analogie avec l'ensemble canonique, nous introduisons une température inverse interne³ m , et une fonction de potentiel :

$$\psi(m) \doteq \frac{1}{N} \log \sum_c 2^{m S_c}, \quad (2.19)$$

supposée automoyennante, et reliée à la complexité $\Sigma(s)$ par une transformation de Legendre :

$$\psi(m) \approx \max_{s \in [s_m, s_M]} [\Sigma(s) + m s] \quad (2.20)$$

Quand ce dernier maximum est atteint à l'intérieur de l'intervalle, on peut écrire :

$$\psi(m) = \Sigma[s^*(m)] + m s^*(m), \quad \text{avec} \quad m = -\partial_s \Sigma[s^*(m)], \quad (2.21)$$

Cette relation sera vérifiée tant que $m \in [m_m, m_M]$, où $s^*(m_m) = s_m$ et $s^*(m_M) = s_M$. La connaissance de $\psi(m)$ permet ainsi de remonter, par transformation inverse, à la complexité $\Sigma(s)$.

La mesure uniforme sur l'ensemble des solutions est en principe décrite par $m = 1$. Cependant, deux cas de figure peuvent se présenter suivant la valeur de m_M (cf. figure 2.4) :

³Habituellement appelée paramètre de brisure de symétrie de Parisi.

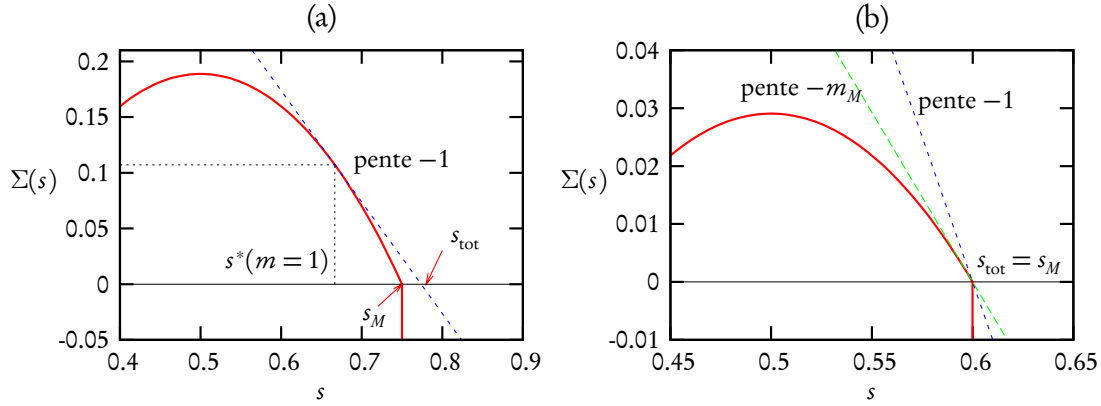


Fig. 2.4: Les deux cas de figures décrits dans le texte. Dans le panneau de gauche (a), l'apposition d'une droite de pente -1 à la courbe de complexité donne le point-col $s^*(1)$. À droite (b), cette droite prend appui sur le point frontière s_M . Celui-ci est décrit par une « température interne » plus élevée ($m_M < 1$), déduite de la transformation de Legendre en s_M : $m_M = -\partial_s \Sigma(s_M)$.

- (a) $m_M > 1$. Le maximum de $\Sigma(s) + s$ est atteint à l'intérieur de l'intervalle de définition. La température inverse effective vaut alors $m = 1$, car l'entropie totale est donnée par $s_{\text{tot}} = \psi(m)/m|_{m=1} = \Sigma[s^*(1)] + s^*(1)$. Bien que l'espace des solutions soit fragmenté, la mesure peut être décrite alternativement soit par un « état » thermodynamique unique, soit par une superposition d'un nombre exponentiel d'états distincts, identifiables aux amas⁴. Nous parlerons de phase *liquide fragmentée* ou encore, pour des raisons historiques, de brisure dynamique de la symétrie des répliques. $s^*(1)$ s'interprète comme l'entropie typique de l'amas contenant une solution prise au hasard avec la mesure uniforme (2.9), et $\Sigma[s^*(1)]$ comme le nombre d'amas concentrant cette mesure.
- (b) $m_M < 1$. Le maximum de $\Sigma(s) + s$ est atteint au bord de l'intervalle de définition, en $s = s_M$, où la complexité s'annule, et où sont vérifiées les relations :

$$s_M = \partial_m \psi(m_M), \quad (2.22)$$

$$\Sigma(s_M) = \psi(m_M) - m_M s_M = -m_M^2 \partial_m \left(\frac{\psi(m)}{m} \right) \Big|_{m_M} = 0. \quad (2.23)$$

La température inverse effective vaut $m_M < 1$, car l'entropie totale est donnée par $s_{\text{tot}} = \psi(m_M)/m_M = s_M$. Ce comportement est en tout point similaire au phénomène de condensation décrit dans le contexte des codes aléatoires au paragraphe

⁴Nous reviendrons plus tard (§4.1.3) sur la définition de la notion d'état, et sur sa relation aux techniques de passage de messages.

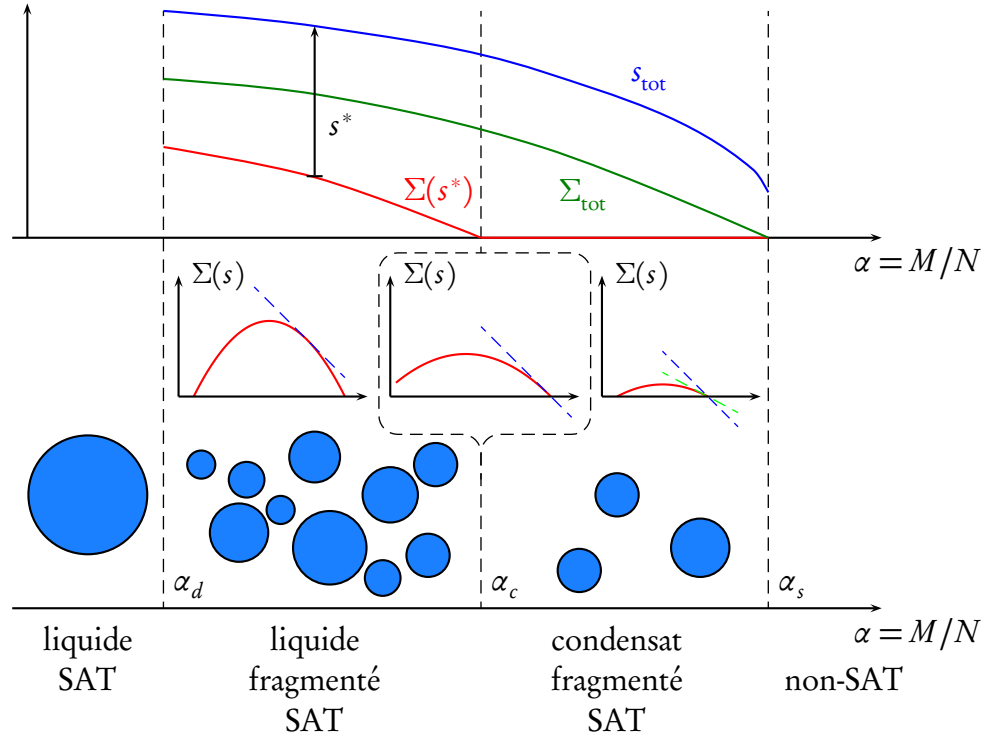


Fig. 2.5: Diagramme complet de k -SAT aléatoire, pour $k \geq 4$. Quand $\alpha < \alpha_d$, la mesure uniforme sur l'ensemble des solutions est dominée par un unique amas. Pour $\alpha_d < \alpha < \alpha_c$, elle est composée d'un nombre exponentiel d'amas, alors que pour $\alpha_c < \alpha < \alpha_s$, ce nombre devient fini. Au point de condensation α_c , la droite de pente -1 prend appui sur la courbe de complexité précisément là où celle-ci s'annule. Dans la partie supérieure du diagramme sont représentées l'entropie totale $s_{\text{tot}} = \Sigma(s^*) + s^*$, la complexité typique $\Sigma(s^*)$ et la complexité totale $\Sigma_{\text{tot}} = \max_s \Sigma(s)$.

	$k = 3$	$k = 4$	$k \rightarrow \infty$
α_d	3,96	9,38	$\frac{2^k}{k} \left[\ln k + \ln \ln k + 1 + O\left(\frac{\ln \ln k}{\ln k}\right) \right]$
α_c	3,96	9,55	$2^k \ln 2 - \frac{3}{2} \ln 2 + O(2^{-k})$
α_s	4,26	9,93	$2^k \ln 2 - \frac{1+\ln 2}{2} + O(2^{-k})$

Tab. 2.1: Seuils de transition dans le problème k -SAT aléatoire, tirés de [KMRT⁺07].

1.2.2. Bien que la température naturelle vaille $m = 1$, l'annulation de la fonction de complexité impose de décrire le système à l'aide d'une température interne plus élevée, appelée température de condensation. On parle alors de phase *fragmentée condensée*, ou de brisure (statique) de la symétrie des répliques. Dans cette phase, la mesure est dominée par la superposition d'un nombre fini d'états thermodynamiques correspondant aux amas d'entropie s_M .

À la lumière de cette classification, il est possible de raffiner le diagramme de phase esquissé figure 2.3. Pour $k \geq 4$, le problème k -SAT aléatoire subit trois transitions de phase alors que la densité $\alpha = M/N$ augmente (voir figure 2.5). En plus des deux transitions déjà évoquées, une transition de *condensation* se produit à l'intérieur de la phase fragmentée, en $\alpha_c \in [\alpha_d, \alpha_s]$, où le nombre d'amas dominant la mesure devient fini. Le nombre *total* d'amas reste néanmoins exponentiel, et est gouverné par la complexité totale $\Sigma_{\text{tot}} = \max_s \Sigma(s)$, qui s'annule en α_s .

Le tableau 2.1 donne quelques valeurs des différents seuils de transition, ainsi que leur comportement asymptotique [MMZ06, KMRT⁺07]. Mentionnons au passage que le cas $k = 3$ échappe au cas général : toute la phase fragmentée y est condensée, de sorte que $\alpha_d = \alpha_c$.

Le phénomène de fragmentation de l'espace des solutions s'accompagne parfois, dans k -SAT aléatoire comme dans d'autres problèmes proches, d'un phénomène de « gel », par lequel certaines variables prennent la même valeur pour toutes les solutions appartenant à un amas donné. Ce comportement est déjà bien connu dans le contexte du problème d'optimisation MAX-SAT, où le fondamental est composé d'un unique amas : l'ensemble des variables gelées est alors désigné par le terme de « colonne vertébrale » [MZK⁺99].

2.3.3 Modèle à amas aléatoires

Nous introduisons un modèle jouet qui généralise le modèle à codes aléatoires du §1.2.2 et reproduit certaines des caractéristiques importantes des problèmes aléatoires de satisfaction de contraintes dans leur phase fragmentée. Au lieu de tirer des mots de codes au hasard, nous tirons $2^{(1-\alpha)N}$ amas au hasard, où α est un paramètre de contrôle.

À chaque amas A nous associons une application aléatoire :

$$\pi_A : \{1, \dots, N\} \longrightarrow \{\{0\}, \{1\}, \{0, 1\}\} \quad (2.24)$$

$$i \longmapsto \pi_A(i) \quad (2.25)$$

telle que pour chaque i , $\pi_A(i)$ vaut $\{0\}$ ou $\{1\}$ avec probabilité $p/2$, et $\{0, 1\}$ avec probabilité $1 - p$. L'amas est alors défini comme suit :

$$A = \{\sigma \mid \forall i \in \{1, \dots, N\}, \sigma_i \in \pi_A(i)\} \quad (2.26)$$

Autrement dit, si $\pi_A(i)$ est un singleton, la variable σ_i est gelée dans l'amas A . Réciproquement, si $\pi_A(i) = \{0, 1\}$, σ_i peut prendre n'importe quelle valeur indépendamment des autres variables dans l'amas A . Notez que le cas $p = 1$ nous ramène à un modèle de codes aléatoires avec $R = 1 - \alpha$.

L'entropie de chaque amas vaut exactement le nombre de variables libres dans cet amas, et suit donc une loi binomiale de paramètres $1 - p$ et N . Ainsi que pour les codes aléatoires, les inégalités de Markov et de Chebychev permettent de montrer que le nombre d'amas d'entropie $S = Ns$ se concentre autour de sa valeur moyenne, de sorte que :

$$\Sigma(s) = \begin{cases} 1 - \alpha - D(s \parallel 1 - p) & \text{p.s. si cette quantité est positive} \\ -\infty & \text{sinon.} \end{cases} \quad (2.27)$$

C'est en fait cette expression de $\Sigma(s)$ que nous avons utilisée pour illustrer le phénomène de condensation dans la figure 2.4, avec $p = 1/2$ et $\alpha = H(1/4)$ à gauche, et $\alpha = H(2/5)$ à droite.

Dans la phase liquide fragmentée, le col $m = 1$ est atteint à l'intérieur de l'intervalle de définition de $\Sigma(s)$, en $s^* = 1 - p/(2 - p)$. La complexité au col :

$$\Sigma(s^*) = \frac{p}{2 - p} + \log(2 - p) - \alpha, \quad (2.28)$$

s'annule au point de condensation $\alpha_c = p/(2 - p) + \log(2 - p)$, tandis que l'entropie totale vaut :

$$s_{\text{tot}} = 1 - \alpha + \log(2 - p). \quad (2.29)$$

Dans le condensat ($\alpha > \alpha_c$), on a $s_{\text{tot}} = s_M$, avec : $D(s_M \parallel p) = 1 - \alpha$. La complexité totale s'écrit quant à elle $\Sigma_{\text{tot}} = 1 - \alpha$ tout le long du diagramme. La figure 2.6 représente ces quantités en fonction de α , ainsi que la température inverse m décrivant le comportement de la mesure.

Nous reprendrons ce modèle jouet dans le chapitre 5 afin d'illustrer certaines propriétés de distances des problèmes de satisfaction de contrainte. Nous verrons également que ce modèle peut s'obtenir, au prix de changements d'échelle adéquats, comme la limite du modèle k -SAT aléatoire à grand k , et du modèle de coloriage sur graphe aléatoire à grand q .

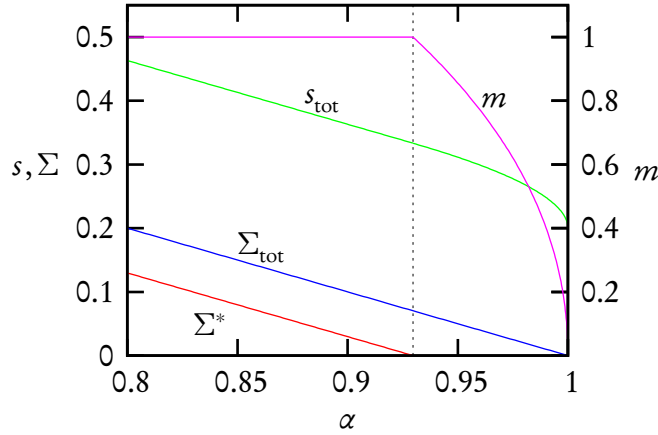


Fig. 2.6: Complexité d'équilibre $\Sigma^* = \Sigma(s^*)$, complexité totale Σ_{tot} et entropie totale s_{tot} du modèle à amas aléatoires en fonction du paramètre α , pour $p = 0.8$.

2.3.4 Ergodicité

Les phénomènes de fragmentation et de condensation entretiennent des liens étroits avec les performances algorithmiques sur les instances aléatoires. En particulier, il a été conjecturé que le phénomène de fragmentation a pour conséquence l'incapacité des algorithmes complets à trouver une solution en temps polynomial. En plus des amas de solutions, il existe des amas « métastables » situés sur des plateaux d'énergie non-nulle. À cause de la distance extensive qui sépare ces amas dans la phase fragmentée, toute dynamique locale à température nulle (basée sur le changement, à chaque pas, d'un nombre fini de variables) est vouée à rester piégée dans le même amas. Par extension, même à température finie, les algorithmes locaux vérifiant le bilan détaillé se trouvent confrontés à des barrières d'énergie de grande taille, et sont condamnés à vivre dans des amas non-optimaux.

Il a été prouvé que tous les algorithmes DPLL basés sur l'heuristique de la clause unitaire trouvent des solutions tant que $\alpha < c_k 2^k / k$, où c_k tend à grand k vers une constante c dépendant de la règle précise de l'algorithme. Parallèlement, le seuil de fragmentation α_d se comporte comme $2^k \ln k / k$ quand k tend vers l'infini. Si l'on met de côté le facteur $\ln k$, qui est par ailleurs cohérent avec la non-universalité la constance c , la similitude entre ces deux comportements est frappante, et milite en faveur de l'interprétation proposée, bien que les algorithmes DPLL soient fort différents des algorithmes de type Metropolis supposés gouverner la dynamique des verres.

Anecdotiquement, l'analogie avec la physique des verres se reflète dans l'usage du terme « complexité » ou « entropie configurationnelle » pour décrire le nombre d'amas. En physique du verre [BB04], l'entropie totale d'un liquide surfondu, c'est-

à-dire l'ensemble de ses degrés de liberté, est la somme d'une entropie vibrationnelle, vivant au sein d'une vallée d'énergie, et d'une entropie configurationnelle correspondant à l'évolution d'une vallée à l'autre. Ces deux entropies sont respectivement s^* et $\Sigma(s^*)$ dans notre langage.

Dans le cas général, une analyse physique des performances algorithmiques a été proposée dans [CMMS04, SM03, SM04], et approfondie dans le cas particulier de XORSAT⁵ dans [MS05, MS06b].

Le type d'approche dont nous rendons compte ici tente d'expliquer des comportements algorithmiques, donc dynamiques, par l'état statique de l'espace configurationnel. Selon ce point de vue, certaines propriétés algorithmiques génériques peuvent être déduites d'une analyse structurelle du problème, indépendamment de l'algorithme considéré. Il faudrait y objecter que ce type de discussion n'est probablement valable que pour une certaine classe d'algorithmes : en effet les expériences numériques montrent que les algorithmes de type *Random WalkSAT* restent performants bien au delà de α_d , battant en brèche l'hypothèse selon laquelle la fragmentation constituerait à elle seule une signature universelle à la performance algorithmique.

Références

Une introduction à la complexité algorithmique et au problème de satisfaisabilité peut être trouvée dans [Pap94]. Du côté de la complexité typique, les premières indications de l'existence d'une transition de phase dans les ensembles aléatoires de problèmes *NP*-complets remontent à [CKT91, SML96]. Cette existence est précisée dans [KS94], tandis que sa relation à la difficulté algorithmique est étudiée dans [MZK⁺99]. L'excellent article de vulgarisation de Bryan Hayes [Hay97] résume certains de ces résultats. Le scénario de fragmentation dans *k*-SAT a été proposé dans [BMW00], et étudié par [MPZ02, MZ02]. Sa présentation sert ici d'introduction aux articles [MMZ05a, MMZ05b], qui en prouvent la validité. La formulation thermodynamique de la statistique des amas est commune à celle adoptée dans les articles [MPR05, MM06b], quoique dans un contexte légèrement différent. Les résultats sur le phénomène de condensation dans *k*-SAT sont tirés de [KMRT⁺07]. Le modèle à amas aléatoires proposé pour illustrer la condensation a d'abord été suggéré par Dimitris Achlioptas, mais introduit indépendamment à l'occasion de la rédaction de cette thèse.

⁵dont la définition est donnée au chapitre suivant.

Chapitre 3

Modèles graphiques

Nous introduisons ici une classe très générale de modèles, définis sur des graphes ou des hypergraphes. Cette présentation, qui servira de cadre aux parties futures, est aussi l'occasion d'introduire les systèmes d'équations linéaires booléennes qui, en dépit de leur apparente simplicité, sont centraux en théorie de l'information et en complexité algorithmique.

3.1 Graphes et hypergraphes

3.1.1 Graphes aléatoires

Un graphe est défini par un ensemble de sommets et par l'ensemble des arêtes les reliant. Un *ensemble aléatoire* de graphes correspond formellement à une mesure sur l'espace de tous les graphes possibles. L'ensemble d'Erdős-Rényi [ER59] est le plus simple d'entre eux : étant donnés N sommets, chacune des $\binom{N}{2}$ arêtes possibles est présente avec une probabilité p . Un autre ensemble fréquemment rencontré est l'ensemble des graphes ℓ -réguliers : dans cet ensemble, sont choisis avec une probabilité uniforme tous les graphes dont chacun des nœuds a exactement ℓ voisins. Ces derniers graphes sont *dilués*, au sens où chaque sommet conserve un nombre fini de voisins quand la taille du graphe tend vers l'infini. La version diluée de l'ensemble d'Erdős-Rényi s'obtient en choisissant $p = \alpha/N$, et génère, quand N tend vers l'infini, une distribution de degrés poissonnienne. En effet, le nombre d'arêtes attachées à un sommet donné suit une loi binomiale :

$$L(\ell) = \binom{N-1}{\ell} \left(\frac{\alpha}{N}\right)^\ell \left(1 - \frac{\alpha}{N}\right)^{N-1-\ell} \rightarrow e^{-\alpha} \frac{\alpha^\ell}{\ell!} \quad (3.1)$$

Il existe un grand nombre de recettes pour contruire des graphes aléatoires dilués.

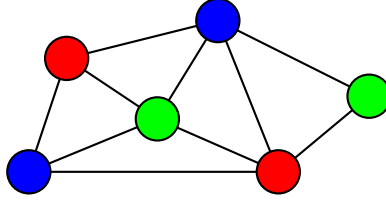


Fig. 3.1: Un exemple de coloriage de graphe.

Une manière assez générale, qui englobe notamment les deux ensembles décrits ci-dessus, consiste à prescrire la distribution des degrés des sommets $L(\ell)$. Dans ce cadre, les graphes ℓ -réguliers sont décrits par $L = \delta_\ell$, et les graphes d'Erdős-Rényi par $L = \text{Poisson}(\alpha)$.

De nombreux problèmes issus de la théorie des probabilités ou de la physique sont définis sur des graphes, comme la percolation sur réseau ou le modèle d'Ising. La transposition d'un problème physique doté d'une dimensionnalité naturelle sur un graphe aléatoire se fait souvent au prix d'une approximation de champ moyen, et laisse parfois de côté des propriétés importantes du système. En retour, elle peut contribuer à la tractabilité analytique du problème.

D'un autre côté, les graphes aléatoires occupent une place importante en théorie des probabilités depuis les travaux fondateurs de Erdős et Rényi [ER59, ER60], et leur utilisation dans l'analyse des réseaux sociaux ou informatiques constitue un important champ d'applications [New03].

3.1.2 Coloriage

Un des problèmes sur graphe les plus étudiés est le problème du coloriage – voir figure 3.1. Un nombre limité q de couleurs étant disponible, il s'agit de colorier les nœuds du graphe de telle sorte qu'aucune paire de sommets voisins ne partagent la même couleur. Ce problème peut être décrit par la fonction de coût suivante :

$$E(\sigma) = \sum_{(i,j) \in A} \left(1 - \delta_{\sigma_i, \sigma_j}\right) \quad (3.2)$$

A désigne ici l'ensemble des arêtes du graphe, et $\sigma_i = 1, \dots, q$ la couleur assignée au sommet i . Comme dans le cas du problème de satisfaisabilité, cette fonction de coût compte le nombre de contraintes violées et les solutions sont les configurations d'énergie nulle.

Pour $q \geq 3$, on sait que le problème de coloriage est *NP*-complet [MT72], tandis qu'il est polynomial pour $q = 2$. En outre, il est remarquable que la version aléatoire

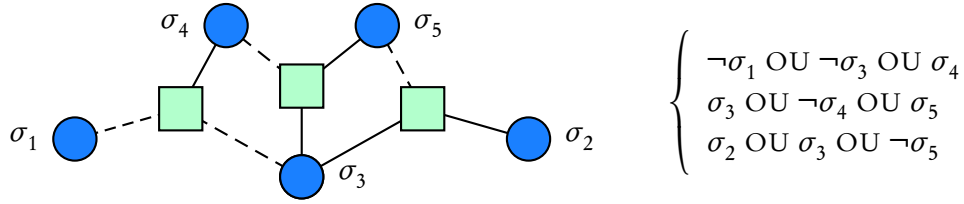


Fig. 3.2: Une instance de 3-SAT représentée par un graphe. Les variables sont représentées par des cercles, et les clauses par des carrés. Chaque clause est la conjonction des variables (ou de leur négation) auxquelles elle est connectée. Une ligne pointillée indique une négation.

de ce problème, définie sur l'ensemble d'Erdős-Rényi dilué avec $p = \alpha/N$, affiche un comportement très proche de celui de la satisfaisabilité dès que $q \geq 3$, avec α dans le rôle de la densité de contraintes. Quand ce paramètre augmente, le système subit dans la limite $N \rightarrow \infty$ une transition coloriable/non-coloriable [AF99] identique à la transition SAT/Non-SAT de la satisfaisabilité, et l'analyse par la physique statistique prévoit également des transitions de fragmentation [MPWZ02] et de condensation [KMRT⁺07, ZK].

3.1.3 Graphes factoriels

Par la nature même des liens du graphe, qui unissent les sommets par deux, les modèles y étant définis reposent sur des interactions à deux corps. Une généralisation naturelle consiste à considérer des *hypergraphes*, ou graphes factoriels, composés de deux types de nœuds : d'une part les sommets où siègent les variables, et d'autre part les nœuds factoriels, ou hyperarêtes, qui relient les variables par une interaction ou une contrainte. Les modèles définis sur de telles structures sont qualifiés de *modèles graphiques*. Le problème de satisfaisabilité se prête bien à une telle description, ainsi que l'illustre la figure 3.2.

Un graphe factoriel est formellement défini par ses nœuds-variables $i \in \{1, \dots, N\}$, ses nœuds factoriels $a \in \{1, \dots, M\}$, et les liens reliant les seconds aux premiers. Étant donné un graphe factoriel, on dotera l'espace configurationnel d'une mesure prenant la forme générale :

$$p(\boldsymbol{\sigma}) = \frac{1}{Z} \prod_{a=1}^M \chi_a(\boldsymbol{\sigma}_a) \quad (3.3)$$

où $\boldsymbol{\sigma}_a$ désigne la collection des variables connectées au facteur a : $\boldsymbol{\sigma}_a = (\sigma_i)_{i \in \partial_a}$. La fonction χ_a peut prendre n'importe quelle forme : par exemple $\chi_a(\boldsymbol{\sigma}_a) = 2^{-\beta E_a(\boldsymbol{\sigma}_a)}$ dans le cas d'une contribution énergétique au poids de Boltzmann, ou encore $\chi_a(\boldsymbol{\sigma}_a) = \mathbb{I}(\boldsymbol{\sigma} \models a)$ pour un problème de satisfaction de contraintes (a désigne alors une clause).

L'ensemble k -SAT aléatoire produit des formules ayant pour hypergraphe sous-jacent un hypergraphe aléatoire tiré selon le même principe que les graphes d'Erdős-Rényi dilués. Notamment, la distribution des degrés ℓ , c'est-à-dire le nombre de clauses auxquelles participe une variable donnée, suit une loi de Poisson :

$$L(\ell) \sim \binom{\binom{N-1}{k-1}}{\ell} \left(\frac{M}{\binom{N}{k}} \right)^\ell \left(1 - \frac{M}{\binom{N}{k}} \right)^{\binom{N-1}{k-1} - \ell} \rightarrow e^{-k\alpha} \frac{(k\alpha)^\ell}{\ell!} \quad (3.4)$$

Cette unité de représentation des modèles graphiques dilués rend moins surprenante la parenté de comportement entre certains problèmes sur graphes aléatoires, comme le coloriage, et d'autres problèmes sur *hypergraphe* aléatoire, comme la satisfaisabilité.

Il existe des alternatives à l'ensemble « poissonnien » que représente k -SAT aléatoire. De la même manière qu'un ensemble de graphes aléatoires peut être défini par sa distribution de degrés $L(\ell)$, une construction générale de graphes factoriels est caractérisée par ses distributions de degrés de variables et de facteurs, respectivement dénotées $L(\ell)$ et $R(k)$. L'ensemble k -SAT aléatoire réalise alors le cas particulier :

$$L = \text{Poisson}(k\alpha) \quad R = \delta_k. \quad (3.5)$$

Des distributions de degrés L et R , on déduit deux autres distributions qui nous seront utiles dans l'exposé de la méthode de cavité au chapitre 4. Supposons que l'on choisisse, au hasard et uniformément, un lien entre une variable et un facteur. On s'intéresse à la distribution des degrés de cette variable et de ce facteur. Le nombre total de liens connectant une variable de degré $\ell + 1$ à un facteur s'évalue à $N(\ell + 1)L(\ell + 1)$. La probabilité de tirer un lien attaché à une variable de degré $\ell + 1$ vaut donc :

$$\lambda(\ell) \doteq \frac{(\ell + 1)L(\ell + 1)}{\mathbb{E}(\ell)}. \quad (3.6)$$

De manière symétrique, la probabilité que le facteur se trouvant à l'autre extrémité du lien ait degré $k + 1$ vaut :

$$\rho(k) \doteq \frac{(k + 1)R(k + 1)}{\mathbb{E}(k)}. \quad (3.7)$$

Ces deux quantités définissent les distributions de degrés dans une perspective d'arête. Par exemple, dans k -SAT aléatoire, on a $\lambda = \text{Poisson}(k\alpha)$ et $\rho = \delta_{k-1}$.

Il est souvent utile de définir la fonction génératrice de ces lois :

$$\begin{aligned} L(x) &= \sum_{\ell} L(\ell)x^\ell, & R(x) &= \sum_k R(k)x^k, \\ \lambda(x) &= \sum_{\ell} \lambda(\ell)x^\ell, & \rho(x) &= \sum_k \rho(k)x^k. \end{aligned} \quad (3.8)$$

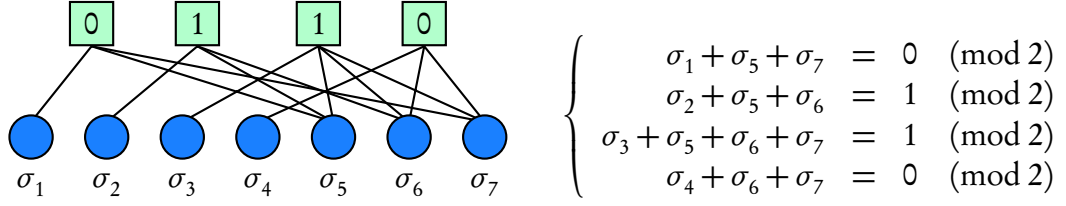


Fig. 3.3: Un système d'équations linéaires booléennes et son graphe de Tanner. Les tests de parité, représentés par des carrés, impliquent les variables qui lui sont adjacentes. Le chiffre à l'intérieur des carrés indique la somme que ces variables doivent prendre modulo 2. Au contraire de la figure 3.2, nous avons choisi de représenter le graphe conformément aux conventions en vigueur en théorie de l'information, avec les variables et les tests sur deux lignes séparées.

En particulier, les connectivités moyennes s'écrivent $\mathbb{E}(\ell) = L'(1)$ et $\mathbb{E}(k) = R'(1)$, et on a :

$$\lambda(x) = \frac{L'(x)}{L'(1)} \quad \rho(x) = \frac{R(x)}{R'(1)}. \quad (3.9)$$

3.2 Équations linéaires booléennes

Nous abordons maintenant un exemple important de problème de satisfaction de contraintes représentable par un graphe factoriel : les systèmes d'équations linéaires booléennes. Ces systèmes, très utilisés en théorie de l'information, sont décrits par un ensemble de tests de parité portant sur des chaînes de bits $\sigma_i \in \{0, 1\}$, de la forme :

$$\sum_{i \in \partial a} \sigma_i = \tau_a \pmod{2}. \quad (3.10)$$

où $\tau_a \in \{0, 1\}$. L'indice a désigne le test de parité, et ∂a l'ensemble des bits présents dans ce test. Graphiquement, les bits σ_i résident sur les sommets indicés par i , et les tests de parité a sont représentés par des hyperarêtes. Dans cette représentation, ∂a est simplement l'ensemble des voisins de a . La figure 3.3 donne un exemple de système linéaire, ainsi que le graphe factoriel correspondant, aussi appelé graphe de Tanner [Tan81] dans ce cas.

Dans la suite nous nous intéressons presque exclusivement aux grandes constructions *aléatoires* d'ensembles de tests de parité.

3.2.1 Le problème XORSAT aléatoire

Suivant le contexte, différents ensembles de systèmes linéaires dilués peuvent être introduits. Nous commençons par décrire le plus simple d'entre eux, l'ensemble k -XORSAT aléatoire [Sch78]. Sur une chaîne de N bits, on tire au hasard M tests de parité impliquant chacun k bits. La valeur de $\tau_a = 0$ ou 1 est tirée avec probabilité $1/2$. Cet ensemble ressemble beaucoup à l'ensemble k -SAT aléatoire, à cette différence près que les clauses sont ici des tests de parité, en lieu et place des fonctions OU. Il paraît par conséquent naturel de définir la limite thermodynamique de la même manière que dans k -SAT, en maintenant le nombre de tests par variable $\alpha = M/N$ constant quand N et M tendent vers l'infini.

Bien que, en tant que système linéaire sur le corps à deux éléments $\mathbb{F}_2$, le problème k -XORSAT soit polynomial (il peut être résolu par élimination de Gauss) il partage avec la satisfaisabilité et le coloriage un diagramme de phases très semblable. La phase fragmentée peut être complètement décrite, et les amas précisément caractérisés, grâce à l'algorithme d'effeuillage que nous décrivons ici brièvement [CDMM03, MRTZ03, MM06b].

Considérons le graphe factoriel de notre problème (cf. figure 3.4), et repérons une variable-sommet connectée à un unique test de parité, appelée « feuille ». Cette variable n'étant pas contrainte par ailleurs, elle assure que le test auquel elle participe pourra toujours être satisfait en ajustant sa valeur. Le problème peut donc être simplifié par l'élimination de la feuille ainsi que de son test de parité. L'algorithme répète cette opération jusqu'à ce qu'il n'y ait plus de feuilles.

Le graphe résiduel obtenu à l'issue de ce processus s'appelle le *cœur*, ou la colonne vertébrale, du problème. Il existe des solutions au problème entier si et seulement s'il en existe au cœur, en vertu de l'argument qui a justifié l'effeuillage. Pour chaque solution du cœur, il existe même un grand nombre (exponentiel en N) de solutions : en effet, lors de l'effeuillage, il peut arriver que plusieurs feuilles soient impliquées dans le même test de parité, créant un degré de liberté par feuille supplémentaire (voir figure 3.5). L'ensemble des solutions associées à une solution de cœur donnée définit un *amas*. Nous justifierons plus tard (§5.2.1) la pertinence de cette définition en la mettant en rapport avec celle que nous avons proposée au chapitre précédent. Remarquons en passant que d'après cette définition, tous les amas ont le même nombre de solutions, et donc la même entropie interne : en effet, le nombre de degrés de libertés gagnés lors de la reconstruction ne dépend pas de la solution de cœur de départ. Par conséquent, à l'inverse de k -SAT ou de la q -colorabilité, le modèle k -XORSAT ne connaît pas de transition de condensation.

L'analyse probabiliste de l'algorithme d'effeuillage permet de distinguer trois comportements suivant la valeur de la densité de tests α . Pour $\alpha < \alpha_d$, l'effeuillage réduit presque sûrement le graphe à un cœur vide : il n'existe alors qu'un seul amas, et

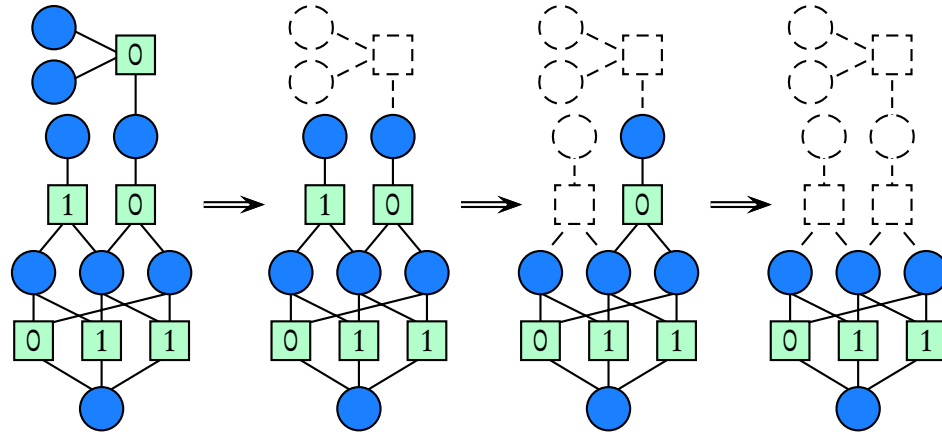


Fig. 3.4: Le processus d'effeuillage. Quand un test de parité compte parmi ses voisins un nœud qui n'est relié à aucun autre test, on peut le supprimer ainsi que tous ses voisins uniquement connectés. Cette opération est répétée jusqu'à ce que toutes les variables soient impliquées dans au moins deux tests de parité. Le graphe restant s'appelle le cœur.

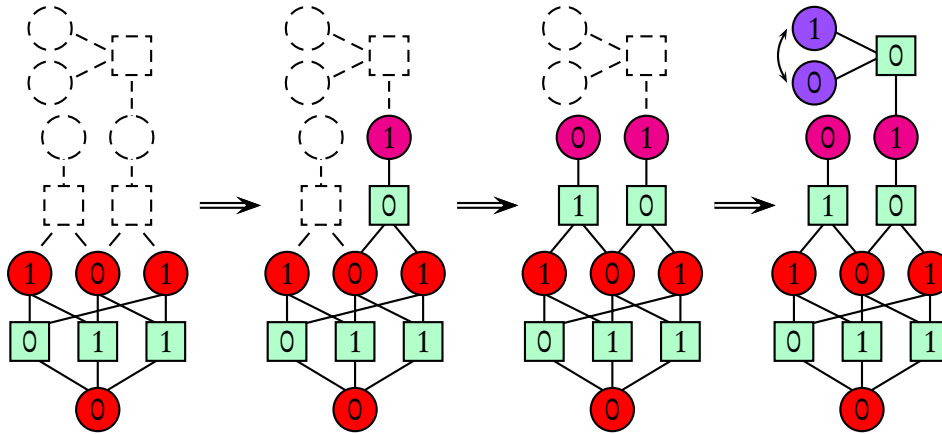


Fig. 3.5: Le processus de reconstruction. À partir d'une solution du cœur, on construit une solution globale en suivant le fil inverse du processus d'effeuillage. À chaque étape, on ajoute un test de parité, ainsi ses variables, auxquelles on assigne une valeur compatible avec le test. Tant que ce choix est unique, et entièrement déterminé par la solution du cœur, on parle de variables gelées. Sinon, les variables sont dites libres. C'est le cas par exemple des deux variables ajoutées lors de la dernière étape, qui peuvent prendre alternativement les valeurs jointes $(0, 1)$ ou $(1, 0)$.

	$k = 3$	$k = 4$	$k = 5$	$k \rightarrow \infty$
α_d	0,818469	0,772278	0,701780	$\sim \ln(k)/k$
α_c	0,917935	0,976770	0,992438	$1 - e^{-k} + O(k^2 e^{-2k})$

Tab. 3.1: Seuils de transition dans le problème k -XORSAT aléatoire.

le problème est résolu en temps linéaire. Quand $\alpha_d < \alpha < \alpha_c$, le cœur garde presque sûrement une taille extensive ; il existe alors un nombre exponentiel de solutions de cœur, et par conséquent un nombre exponentiel d'amas. Quand $\alpha > \alpha_c$ ($\alpha_c < 1$), le cœur devient sur-contraint, et plus aucune solution n'existe : cette seconde transition est la transition SAT/non-SAT habituelle. La table 3.1 reproduit les seuils α_d et α_c pour diverses valeurs de k .

Tout le long de la phase SAT, le logarithme du nombre de solutions vaut $N - M$ presque sûrement, et subit une brusque transition vers $-\infty$ en α_c . En termes d'algèbre linéaire, cela se traduit par le fait que les tests de parité sont typiquement indépendants dans la phase SAT. Si tel n'était pas le cas, une contradiction se produirait avec probabilité $1/2$ pour chaque combinaison linéaire les liant. La discontinuité en α_c , bien que surprenante au premier abord, est en fait corrélative de l'existence d'une phase fragmentée. À la transition SAT/non-SAT, c'est en effet le logarithme du nombre d'amas qui s'annule continûment.

L'algorithme de reconstruction (Fig. 3.5) permet d'identifier deux types de variables : d'un côté, les variables « gelées » prennent la même valeur dans toutes les solutions d'un même amas, soit parce qu'elles appartiennent au cœur, soit parce qu'elles se déduisent directement de celui-ci ; de l'autre côté, les variables « libres », qui peuvent prendre les valeurs 0 ou 1. Nous avons indiqué au chapitre précédent que ce phénomène de gel était caractéristique des phases difficiles des problèmes de satisfaction de contraintes, et nous nous sommes même appuyés dessus afin de construire le modèle jouet d'amas aléatoires. Le cas de XORSAT offre néanmoins une simplification majeure : l'ensemble des variables gelées *ne dépend pas de l'amas considéré*, et se déduit simplement de l'algorithme d'effeuillage.

Nous reviendrons par la suite sur ce problème qui, en dépit de sa simplicité, présente dans une version purifiée certaines propriétés caractéristiques des problèmes aléatoires difficiles, et permet une exposition clarifiée des concepts et outils y intervenant.

3.2.2 Utilisation pour la compression de données

La version « optimisation » du problème k -XORSAT trouve une application intéressante en théorie de l'information, et plus précisément dans le contexte de la compression de données avec perte évoqué au paragraphe 1.2.3. Mettons que l'on veuille compresser une chaîne de bits τ_a , $a = 1, \dots, M$ en une chaîne σ_i , $i = 1, \dots, N$ plus courte. Pour cela, on cherche la chaîne σ qui satisfasse le maximum de tests de parité de la forme (3.10) (codage). La chaîne τ^* reconstituée à partir de σ s'obtient par ces mêmes équations (décodage) :

$$\tau_a^* = \sum_{i \in \partial a} \sigma_i \pmod{2} \quad (3.11)$$

et la distorsion, c'est-à-dire le nombre de bits a tels que $\tau_a^* \neq \tau_a$, est donnée par le nombre de tests violés dans la recherche de σ , autrement dit l'énergie fondamentale du problème d'optimisation¹. Celle-ci peut être calculée par les outils de la physique statistique [CM05], et il a été montré qu'elle sature rapidement la borne de Shannon quand k grandit [MW06]. Nous montrons ici ce résultat directement en considérant un code linéaire parfaitement aléatoire :

$$\sum_{i=1}^N \lambda_{ai} \sigma_i = \tau_a \pmod{2} \quad (3.12)$$

où λ_{ai} vaut 0 ou 1 avec probabilité 1/2. Ce modèle est supposé correspondre à la limite de connectivité infinie $k \rightarrow \infty$. De la même manière que le modèle à énergies aléatoires (REM) s'obtient comme la limite $p \rightarrow \infty$ du modèle p -spin, on s'attend à retrouver ici les caractéristiques d'un modèle à mots de codes aléatoires tel que celui décrit dans la section 1.2.3. Afin de confirmer cette intuition, il est pratique de reformuler le problème en termes d'algèbre linéaire. La distorsion est alors interprétée comme la distance entre la chaîne à compresser τ et l'image de la matrice $\Lambda = \{\lambda_{ai}\}$:

$$D(\tau) = \min_{\sigma} \|\Lambda \sigma - \tau\| \quad (3.13)$$

En effet, l'ensemble des mots de code est ici l'image d'une matrice booléenne aléatoire. Les propriétés de distance de cet ensemble diffèrent peu de celui des mots de codes complètement aléatoires. En particulier, un rapide calcul combinatoire montre que le nombre de mots de codes à distance donnée w de τ a pour moyenne et pour variance les mêmes valeurs que dans le modèle à codes aléatoires :

$$\mathbb{E}(n_w) \doteq \mathbb{E} \text{card}\{\sigma \mid \|\Lambda \sigma - \tau\| = w\} = \binom{M}{w} 2^{N-M} \asymp 2^{M[H(w/M) + \alpha^{-1} - 1]}, \quad (3.14)$$

¹Comme au §1.2.3, c'est le décodage qui constitue ici la partie difficile, tandis que le codage est trivial.

$$\text{Var } n_w = \mathbb{E}(n_w) \left[1 - 2^{-M} \binom{M}{w} \right]. \quad (3.15)$$

Ainsi, en vertu des mêmes arguments que ceux utilisés pour les codes aléatoires dans la section 1.2, le taux de distorsion D/M tend presque sûrement vers la borne de Shannon $\delta_{GV}(R)$, où $R = \alpha^{-1}$ est le taux de compression.

Bien que la performance du code reste théoriquement satisfaisante à k fini, le problème d'optimisation posé par le codage présentent des difficultés, qu'il est possible de contourner à l'aide de généralisations astucieuses de l'algorithme de propagation des sondages (*survey propagation*, cf. chapitre 6) [WM03]. Une version modifiée, où les tests de parité sont remplacés par des portes logiques aléatoires, se prête en revanche directement à la propagation des sondages, permettant une compression quasi-optimale en un temps raisonnable [CMZ05a, CMZ05b]. Un autre ensemble de problèmes linéaires [MO03] semble également pouvoir se résoudre à l'aide d'algorithmes de passage de messages [Mur04].

3.2.3 Les codes linéaires dilués

Nous passons maintenant à une autre application, beaucoup plus développée, des codes linéaires en théorie de l'information, à savoir leur utilisation comme codes de correction d'erreur. Les codes à faible densité de tests de parité (*low-density parity checks codes*, LDPC), introduits par Gallager [Gal62, Gal68], utilisent l'espace vectoriel des solutions d'un problème linéaire comme livre de code. Les mots de code sont donc les solutions de M équations du type (3.10), avec $\tau_a = 0$. Quand la matrice a rang maximal, l'espace des mots de code a pour dimension $L = N - M$, ce qui signifie que chaque mot transporte L bits d'informations. Le taux du code est le rapport de la quantité d'information codée sur la longueur des mots de code : $R = L/N = 1 - M/N$.

Ce schéma se distingue de celui de la compression par plusieurs aspects : dans la compression, c'est l'*image* d'une matrice booléenne aléatoire qui définit l'ensemble des mots de code, alors que pour la correction d'erreur c'est le *noyau* qui joue ce rôle. Par ailleurs, le décodage par mots constitue dans le cas présent la partie la plus difficile : il s'agit, rappelons-le, de trouver le mot de code le plus proche du message reçu. Ce problème d'optimisation est en général très difficile, et le problème de décision associé est NP-complet. En pratique, il peut être résolu de manière sous-optimale par des méthodes d'échange de message, que nous décrirons au chapitre suivant.

Les constructions aléatoires, quand elles sont utilisées dans le contexte de la correction d'erreurs, doivent remplir certaines conditions précises. Par exemple, les codes où certains bits ne sont contraints que par un — ou aucun — test de parité, comme c'est le cas dans le problème k -XORSAT, sont à exclure, car de tels bits ne sauraient être corrigés en cas de corruption (si non contraints) ou ne contiendraient pas d'information

(si contraints par un seul test). Un ensemble classique et très étudié de codes LDPC est l'ensemble régulier (ℓ, k) , où chaque bit est contraint par ℓ tests, et où chaque test comporte k bits. Cependant, il est possible d'utiliser une large fourchette de degrés de bits et de tests, en construisant des graphes irréguliers aléatoires caractérisés par leurs distributions de degrés $L(\ell)$ et $R(k)$. La condition de connectivité $M\mathbb{E}(k) = N\mathbb{E}(\ell)$ permet d'exprimer le taux du code en fonction des degrés moyens :

$$R = 1 - \frac{\mathbb{E}(\ell)}{\mathbb{E}(k)}. \quad (3.16)$$

Il s'avère en fait que les constructions irrégulières sont les plus efficaces du point de vue des performances algorithmiques. Nous illustrerons plus tard cette assertion dans le contexte du canal d'effacement.

Dans la limite des grands mots ($N \rightarrow \infty$), alors qu'on augmente le paramètre de corruption ϵ du canal (BSC ou BEC), les constructions aléatoires de codes LDPC, régulières comme irrégulières, subissent une transition abrupte, d'un régime où presque tous les messages transmis peuvent être décodés sans erreur, vers un régime où presque aucun ne peut l'être. Le bruit critique ϵ_c pour lequel cette transition se produit est toujours inférieur à la borne de Shannon.

Limite des codes aléatoires

La borne de Shannon peut être saturée en augmentant la connectivité, comme avec les codes de compression. La limite de grande connectivité est modélisée par un code linéaire aléatoire, dont les mots de code vérifient les M équations :

$$\sum_{i=1}^N \lambda_{ai} \sigma_i = 0 \pmod{2} \quad (3.17)$$

où λ_{ai} vaut 0 ou 1 avec probabilité 1/2.

Détaillons les performance de ce code pour le canal symétrique BSC. Notons tout d'abord que la structure en groupe de l'ensemble des mots de code permet de se ramener au mot de code $\mathbf{0} = (0, \dots, 0)$ sans perte de généralité. Quand ce mot passe dans le canal symétrique, chaque bit est changé en 1 avec probabilité ϵ . Si le décodage utilise le principe de vraisemblance maximale par mots, il sera réussi si le mot reçu est plus proche de $\mathbf{0}$ que de n'importe quel autre mot de code. Pour évaluer la probabilité de cet évènement, nous reprenons l'argumentaire de la section 1.2.2 et estimons le nombre n_w de mots de codes (autre que $\mathbf{0}$) en fonction de la distance au mot reçu. Un peu de combinatoire donne :

$$\mathbb{E}(n_w) = 2^{-M} \binom{N}{w} [1 - \epsilon^d (1 - \epsilon)^{N-w}] \asymp 2^{N[H(w/N) + R - 1]}, \quad (3.18)$$

$$\text{Var } n_w = \mathbb{E}(n_w) (1 - 2^{-M}). \quad (3.19)$$

Les inégalités de Markov et de Chebychev montrent ainsi que le plus proche des « mauvais » mots de code est presque sûrement à distance $\sim N\delta_{GV}(R)$. Le mot de code original $\mathbf{0}$ étant à distance $\sim N\epsilon$ presque sûrement, le décodage réussira si et seulement si $\epsilon < \delta_{GV}(R)$. La borne de Shannon est donc bien atteinte.

Le cas du BEC se traite de manière similaire, avec le même résultat. La limite de grands degrés permet ainsi d'approcher le cas des codes aléatoires, qui saturent la borne de Shannon.

Algorithme d'effeuillage, bis

Nous avons jusqu'à maintenant discuté les performances *optimales* des codes LDPC. En pratique cependant, il faut souvent recourir à des algorithmes sous-optimaux. Nous examinons maintenant un algorithme simple [LMS⁺97] permettant de décoder les codes LDPC sur le canal d'effacement. Voici comment il procède :

1. On assigne aux bits correctement reçus leur valeur $\sigma_i = 0$ ou 1 , et aux bits effacés la valeur « joker » $\sigma_i = *$. Les bits connus sont « nettoyés » du graphe factoriel, ainsi que les tests auxquels ne participent que des bits connus.
2. Tant qu'il existe, dans le graphe résiduel, un test de parité a ayant exactement *un* voisin de valeur indéterminée, noté i ($\sigma_i = *$) :
 - Assigner $\sigma_i \leftarrow \sum_{j \in \partial a \setminus i} \sigma_j$.
 - Nettoyer le graphe en supprimant a et i .

Le décodage est un succès si et seulement si, à l'issue de la procédure, il ne reste plus de graphe. Cet algorithme est le dual exact de l'algorithme d'effeuillage proposé pour k -XORSAT. Remarquons que, ici encore, le succès ne dépend pas du mot de code envoyé.

L'étude statistique de l'algorithme sur un ensemble de codes aléatoires prévoit qu'en dessous d'une certaine valeur critique du bruit ϵ_d , l'effeuillage mange le graphe presque sûrement quand $N \rightarrow \infty$. En revanche, au dessus de ce seuil, l'algorithme est presque sûrement bloqué par un graphe résiduel dont tous les facteurs ont deux voisins ou plus. De tels graphes, dont l'ensemble des variables est appelé sous-parties d'arrêt (*stopping sets*), jouent donc un rôle important dans la détermination des performances algorithmiques des codes LDPC sur le canal d'effacement [DPTTJR02]. Ils sont l'équivalent du *cœur* de k -XORSAT.

Pour autant, l'arrêt de l'effeuillage avant la disparition totale du graphe n'implique pas que le graphe résiduel admette plusieurs solutions : l'apparition de solutions parasites est en effet déterminée par le seuil ϵ_c , qui est strictement supérieur à ϵ_d . Quand $\epsilon_d < \epsilon < \epsilon_c$, le système d'équations représenté par la sous-partie d'arrêt à

l'issue de l'algorithme admet bien une solution unique, mais l'algorithme d'effeuillage échoue à la trouver. C'est en ce sens que cet algorithme est *sous-optimal*.

3.3 Problèmes d'occupation

Les sous-parties d'arrêt, que nous venons d'introduire, appartiennent à une classe plus générale de problèmes *d'occupation* définis sur des graphes factoriels. Une sous-partie d'arrêt S est un sous-ensemble des variables d'un graphe factoriel, vérifiant :

$$\text{pour chaque facteur } a, \quad |\partial a \cap S| \neq 1. \quad (3.20)$$

Autrement dit, chaque facteur est relié à au moins deux variables de S , à moins qu'il soit complètement isolé — auquel cas il peut être considéré comme hors du sous-graphe défini par S .

Une généralisation naturelle consiste à chercher dans un graphe des sous-parties avec des propriétés de degré particulières. En général :

$$\text{pour chaque facteur } a, \quad |\partial a \cap S| \in A, \quad \text{où } A \subset \mathbb{N}. \quad (3.21)$$

Les sous-parties S vérifiant cette condition sont appelées A -parties. S peut être décrit par une chaîne binaire σ , où $\sigma_i = 1$ si $i \in S$, et $\sigma_i = 0$ sinon. La condition devient alors : $\forall a, \sum_{i \in \partial a} \sigma_i \in A$. Il s'avère que de nombreux problèmes classiques rentrent dans ce cadre :

- Le problème de couverture de graphe [WH00, WH01] : on «couvre» les nœuds d'un graphe simple en s'assurant que, parmi deux nœuds voisins, au moins l'un d'entre eux est couvert. S désigne l'ensemble des nœuds couverts, et le graphe simple est transformé en graphe factoriel par l'insertion d'un facteur sur chaque lien du graphe. Le problème de couverture revient alors à poser $A = \{1, 2\}$.
- Les cycles d'un graphe [MMS06, MS06a]. On transforme un graphe simple G en graphe factoriel F de la manière suivante : on met une variable sur chaque arête du graphe, et un facteur sur chaque nœud. Une partition de cycles disjoints de G est définie par une collection S d'arêtes (c'est-à-dire de variables dans le graphe factoriel) telle que chaque nœud de G soit relié à exactement à deux arêtes appartenant à S , ou à aucune. Dans notre formalisme, cela se traduit par $A = \{0, 2\}$.
- Les appariements, ou dimères, sur graphe [LP86, ZM06]. Un dimère est une paire de nœuds voisins d'un graphe simple. D'après une règle d'exclusivité, chaque nœud ne peut appartenir qu'à un dimère. La présence d'un dimère est codé par une variable siégeant sur chaque arête. Comme précédemment, on place

- un facteur sur chaque nœud. La condition d'exclusivité équivaut alors à poser $A = \{0, 1\}$.
- Les sous-graphes r -réguliers [PW06]. Basé sur la même transformation que les deux cas précédents, avec $A = \{0, k\}$.
 - La couverture exacte d'un graphe factoriel, aussi appelé 1-parmi- k -SAT positif [KM05, RSZ07]. Tous les facteurs ont degré k , et $A = \{1\}$. C'est le complémentaire exact des sous-parties d'arrêt.
 - Le bicoloriage d'un graphe [CNRTZ03]. Dans un graphe factoriel, σ_i représente l'une des deux couleurs que peut prendre une variable. La contrainte est que tous les voisins d'un facteur n'aient pas la même couleur. Si le facteur a degré k , $A = \{1, \dots, k-1\}$.
 - Les mots de codes d'un LDPC. Chaque test de parité imposant que la somme des variables soit paire, ce cas est décrit par $A = 2\mathbb{N}$. L'énumération des mots de codes en fonction de leur « poids » $|S|$ renseigne sur les propriétés de distance du code, qui sont cruciales pour la compréhension des régimes de faible bruit [DRU06, DMU04].

Les problèmes de percolation et, plus généralement, de χ -cœur [Bol01] relèvent du même type de formulation. Il est intéressant de remarquer que dans presque tous ces problèmes, des algorithmes « d'effeuillage » ont été proposés.

Les sous-parties d'arrêt jouent, dans le décodage itératif par l'algorithme d'effeuillage, le même rôle que les mots de code dans le décodage optimal : une fois le graphe débarrassé des bits reçus, l'effeuillage réussira pour peu que le graphe nettoyé soit exempt de sous-parties d'arrêt non-triviales ($A = \mathbb{N} \setminus \{1\}$) tandis que le décodage optimal ne réussira que si le graphe nettoyé est exempt de mots de codes non-triviaux ($A = 2\mathbb{N}$). L'inclusion stricte $2\mathbb{N} \subsetneq \mathbb{N} \setminus \{1\}$ implique la sous-optimalité du décodage itératif.

Références

Le livre de Bollobás [Bol01] contient une somme importante de problèmes et de résultats intéressants sur les graphes aléatoires. Une autre référence classique en la matière est [JLR00]. La notion de graphe factoriel, qui est au cœur de la plupart des travaux de cette thèse, est introduite dans de nombreux ouvrages et tutoriaux [KFL01, Mac03, RU07, MM07]. Les algorithmes d'effeuillage, ainsi que les notions de cœur ou de sous-partie d'arrêt, ont été découverts indépendamment dans des contextes différents [CDMM03, MRTZ03, LMS⁺97]. Ils s'inscrivent en fait dans le cadre plus large du problème de la recherche du χ -cœur [Bol01, JLR00, PSW96], qui généralise la notion de percolation. Les articles [MR06a, MR06b] et [MM06b] utilisent largement ces algorithmes d'effeuillages, et exploitent leur relation aux tech-

niques de passage de messages.

Les codes linéaires dilués (LDPC) ont été découverts par Gallagher [Gal62, Gal68]. Jugés impraticables en l'état des capacités calculatoires de l'époque, ils ont été laissés en jachère jusqu'à leur « redécouverte » par MacKay et Neal [MN95, MN96, Mac99, LMS⁺97]. Considérés comme les codes les plus performants à ce jour, ils sont l'objet d'un livre de synthèse en préparation [RU07]. Depuis l'identification des codes linéaires à des modèles de spins [Sou89, Sou94], les codes LDPC ont fait l'objet d'une activité soutenue de la part de la physique statistique [Nis01, KS04]. Les codes LDPC constituent la base d'étude des articles [MR06a, MR06b].

Chapitre 4

Passage de messages

Les algorithmes basés sur l'échange de messages ont fait la preuve de leur efficacité dans des domaines aussi variés que la communication, l'inférence, l'optimisation ou la physique statistique. Ce chapitre présente ces méthodes dans un cadre unifié inspiré de la physique, et met l'accent sur les applications en théorie de l'information et en optimisation.

4.1 Approximation des arbres

La méthode de cavité fut introduite originellement dans le contexte des modèles de spins en champ moyen comme une alternative à la méthode des répliques [MPV86, MPV87]. L'étude se limitait alors à des graphes ou hypergraphes *complets*, où toutes les arêtes possibles étaient présentes. L'extension de cette méthode aux modèles graphiques dilués [MP01] s'avère en fait équivalente, dans l'hypothèse de la symétrie des répliques, à l'*approximation de Bethe* [Bet35], qui est exacte sur les arbres. Par ailleurs, pour les graphes dilués comme pour les graphes complets, la méthode de la cavité fournit les mêmes résultats que la méthode des répliques [Mon98] après moyennage sur les instances.

Indépendamment de ces travaux, relatifs à l'étude des systèmes vitreux, des techniques de passage de messages ont été depuis longtemps développées afin de résoudre des problèmes de communication et d'inférence. Gallagher, dans sa thèse de 1962 [Gal62], fut sans doute le premier à proposer une série d'algorithmes basés sur cette idée pour traiter ses codes linéaires dilués. Depuis, de nombreux progrès ont été accomplis, donnant forme à la version la plus efficace et la plus utilisée de ces techniques : l'algorithme de « propagation des convictions » (*Belief Propagation*, BP), aussi connu sous le nom d'algorithme « somme-produit » [KFL01]. Dans les bons cas, cet algorithme converge vers une solution qui réalise précisément l'approximation de Bethe [YFW02], et donc celle de la cavité.

Derrière ces appellations variées se cache donc une même et unique méthode, que nous exposons ici, en commençant par l'exemple simple d'une chaîne linéaire de spins d'Ising.

4.1.1 Chaîne d'Ising

Le problème d'Ising à une dimension est caractérisé par un Hamiltonien de la forme :

$$E(\sigma) = - \sum_i (h_i \sigma_i + J_i \sigma_i \sigma_{i-1}) \quad (4.1)$$

avec pour mesure :

$$p(\sigma) = \frac{1}{Z} 2^{-E(\sigma)}. \quad (4.2)$$

Nous avons déjà vu au paragraphe 1.1.4 que ce modèle peut être facilement résolu, à condition de lui donner une forme markovienne. La méthode de la cavité, bien que relevant d'une idée apparemment distincte, réalise cette transformation.

On définit la marginale de cavité $p_{i \rightarrow i+1}(\sigma_i)$ comme la probabilité que le spin i prenne la valeur σ_i , *une fois coupé le lien entre i et $i + 1$* (voir figure 4.1a), et on lui associe le champ de cavité $h_{i \rightarrow i+1}$:

$$p_{i \rightarrow i+1}(\sigma_i) = \frac{2^{h_{i \rightarrow i+1} \sigma_i}}{2 \cosh h_{i \rightarrow i+1}} \quad (4.3)$$

Ici, à l'instar des logarithmes, les fonctions hyperboliques sont en base 2, et la température inverse β est fixée à 1, sans perte de généralité. Le champ de cavité peut être calculé récursivement à l'aide de la formule :

$$p_{i \rightarrow i+1}(\sigma_i) \propto \sum_{\sigma_{i-1} = \pm 1} p_{i-1 \rightarrow i}(\sigma_{i-1}) 2^{-h_i \sigma_i - J_i \sigma_i \sigma_{i-1}} \quad (4.4)$$

d'où l'on déduit :

$$h_{i \rightarrow i+1} = h_i + \operatorname{arctanh}(\tanh J_i \tanh h_{i-1 \rightarrow i}) \doteq h_i + u_{i-1 \rightarrow i} \quad (4.5)$$

La quantité $u_{i-1 \rightarrow i}$, appelée *biais* de cavité, mesure l'influence de la variable $i - 1$ sur le champ effectif en i . Symétriquement, si l'on dénote par $h_{i \rightarrow i-1}$ le champ de cavité en i quand le lien entre i et $i - 1$ est coupé, on obtient dans l'autre sens :

$$h_{i \rightarrow i-1} = h_i + \operatorname{arctanh}(\tanh J_{i+1} \tanh h_{i+1 \rightarrow i}) \doteq h_i + u_{i+1 \rightarrow i}. \quad (4.6)$$

Le champ effectif appliqué en i , une fois rétablis les liens avec $i - 1$ et $i + 1$, s'obtient en fonction des champs de cavité en $h_{i+1 \rightarrow i}$ et $h_{i-1 \rightarrow i}$ (figure 4.1b). Il est important

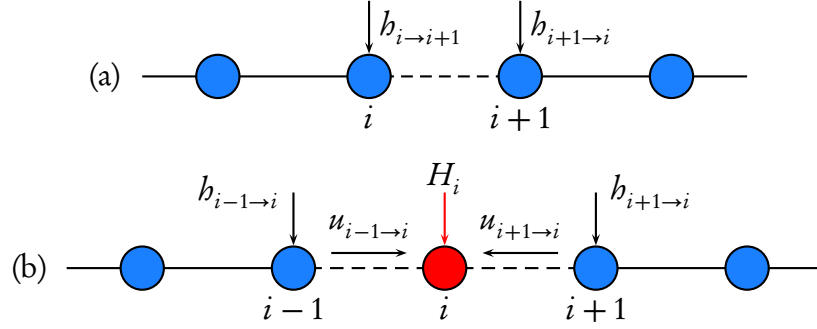


Fig. 4.1: (a) Définition des champs de cavité. Quand le lien entre i et $i+1$ est supprimé, $h_{i \rightarrow i+1}$ et $h_{i+1 \rightarrow i}$ désignent respectivement les champs effectifs en i et $i+1$. (b) Calcul du champ effectif en i . Le rétablissement des liens en pointillés soumet la variable i à l'Hamiltonien effectif local (4.7), induisant un champ local H_i .

de noter que les marginales de cavité correspondant à ces champs sont *indépendantes entre elles* : en effet, les valeurs de spins en $i-1$ et $i+1$ ne sont corrélées qu'en présence du site i . Un calcul sous l'Hamiltonien effectif local

$$E_i = -h_i \sigma_i - h_{i-1 \rightarrow i} \sigma_{i-1} - h_{i+1 \rightarrow i} \sigma_{i+1} - J_i \sigma_i \sigma_{i-1} - J_{i+1} \sigma_{i+1} \sigma_i \quad (4.7)$$

donne donc :

$$H_i \doteq \frac{1}{2} \log \frac{p_i(+1)}{p_i(-1)} = h_i + u_{i-1 \rightarrow i} + u_{i+1 \rightarrow i} \quad (4.8)$$

Cette formule est-elle compatible avec le résultat du paragraphe 1.1.4 ? Rappelons que le taux de transition $q(\sigma|\sigma')$ s'écrit, en convention de spins :

$$q_i(\sigma|\sigma') = 2^{a_i \sigma + b_i \sigma' + c_i \sigma \sigma' + d_i} \quad (4.9)$$

Les deux conditions de normalisation impliquent :

$$b_i = -\operatorname{arctanh}(\tanh a_i \tanh c_i), \quad (4.10)$$

tandis que l'équation (1.37) se traduit par :

$$H_i = a_i + \operatorname{arctanh}[\tanh(H_{i-1} + b_i) \tanh c_i] \quad (4.11)$$

Avec les changements de variables $c_i = J_i$ et $h_i = a_i + b_{i+1}$, et l'identification $b_i = -u_{i \rightarrow i-1}$, $a_i = h_{i \rightarrow i-1}$, on retrouve ainsi les équations de cavité.

Comme dans le cas de la chaîne de Markov, cf. (1.36), la méthode de la cavité s'accompagne d'une manière de calculer l'énergie libre, c'est-à-dire, au signe près, le

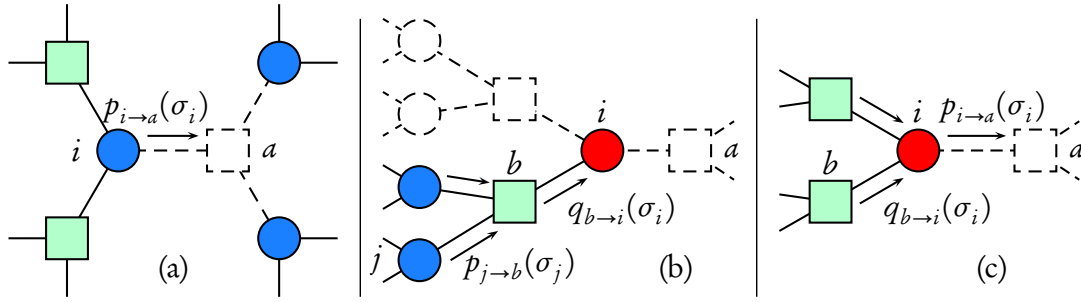


Fig. 4.2: La méthode de la cavité. (a) $p_{i \rightarrow a}$ désigne la loi de σ_i quand le facteur a est absent. (b) Récursion de cavité : on commence par calculer $q_{b \rightarrow i}(\sigma_i)$, qui mesure le poids ressenti par la variable i du fait de la présence de b . Ce poids s'exprime en fonction des marginales de cavité entrantes $p_{j \rightarrow b}$, qui sont factorisées cf. (4.17) : en effet, dans l'hypothèse de l'arbre, deux variables adjacentes à b ne peuvent être corrélées qu'en présence de b . (c) La nouvelle marginale de cavité $p_{i \rightarrow a}(\sigma_i)$ est proportionnelle au produit des poids $q_{b \rightarrow i}(\sigma_i)$ cf. (4.18). Là encore l'indépendance est utilisée pour justifier la factorisation : les sous-arbres attachés aux voisins de i ne sont corrélés qu'en présence de i .

logarithme de la constante de normalisation. Chaque site i contribue à cette quantité par l'Hamiltonien local (4.7), d'où une contribution :

$$-F_i = \log(2 \cosh H_i). \quad (4.12)$$

Cependant, l'énergie libre ne saurait être simplement la somme de ces contributions, car chaque lien serait alors compté deux fois. Il convient donc de soustraire une fois chaque contribution de lien. Un lien $(i, i+1)$ étant soumis à l'Hamiltonien local

$$E_{i,i+1} = -h_{i \rightarrow i+1} \sigma_i - h_{i+1 \rightarrow i} \sigma_{i+1} - J_{i+1} \sigma_i \sigma_{i+1} \quad (4.13)$$

cette contribution s'écrit :

$$-F_{i,i+1} = \log(1 + \tanh J_{i+1} \tanh h_{i+1 \rightarrow i} \tanh h_{i \rightarrow i+1}). \quad (4.14)$$

L'énergie libre vaut donc :

$$F = -\log Z = \sum_i (F_i - F_{i,i+1}). \quad (4.15)$$

Il est facile de vérifier que cette formule est compatible avec celle obtenue équation (1.36).

4.1.2 Ramification de branches

L'étude précédente, qui résout exactement le problème de la chaîne d'Ising, peut être généralisée à une mesure de la forme (3.3) :

$$p(\sigma) = \frac{1}{Z} \prod_{a=1}^M \chi_a(\sigma_a), \quad (4.16)$$

pourvu que le graphe sous-jacent soit un arbre. On définit la marginale de cavité $p_{i \rightarrow a}(\sigma_i)$, où $i \in \partial a$, comme la loi de probabilité de la variable σ_i en l'absence du facteur a (voir figure 4.2a). Les équations de cavité prennent alors la forme :

$$p_{i \rightarrow a}(\sigma_i) = \hat{p}[\{q_{b \rightarrow i}\}_{b \in \partial i \setminus a}](\sigma_i) \doteq 2^{F_{i \rightarrow a}} \prod_{b \in \partial i \setminus a} q_{b \rightarrow i}(\sigma_i), \quad (4.17)$$

$$q_{b \rightarrow i}(\sigma_i) = \hat{q}[\chi_b, \{p_{j \rightarrow b}\}_{j \in \partial b \setminus i}](\sigma_i) \doteq \sum_{\sigma_{b \setminus i}} \chi_b(\sigma_b) \prod_{j \in \partial b \setminus i} p_{j \rightarrow b}(\sigma_j) \quad (4.18)$$

où $\sigma_{a \setminus i}$ désigne la collection des variables σ_j , pour $j \in \partial a \setminus i$. Ces équations donnent la nouvelle marginale de cavité $p_{i \rightarrow a}$ à l'issue de la ramification en i des variables $j \in \partial^2 i \setminus a$ (les seconds voisins de i ne transitant pas par a), par l'intermédiaire des facteurs $b \in \partial i \setminus a$ (cf. figure 4.2b et 4.2c). Ces variables sont, avant cette ramification, indépendantes entre elles, ce qui autorise la factorisation de leur loi jointe. Les variables $q_{a \rightarrow i}$, qui servent de quantités intermédiaires, admettent une interprétation simple (cf. figure 4.2b) : $q_{a \rightarrow i}(\sigma_i)$ est proportionnel à la marginale de σ_i dans l'hypothèse où le site i n'est connecté qu'au facteur a . Finalement, le facteur de normalisation $2^{-F_{i \rightarrow a}}$ mesure la variation d'énergie libre consécutive au branchement en i des sous-arbres de racine j .

Le calcul des vraies marginales, sans suppression de facteur, s'effectue de manière identique :

$$p_i(\sigma_i) = 2^{F_{i+a \in \partial i}} \prod_{a \in \partial i} q_{a \rightarrow i}(\sigma_i). \quad (4.19)$$

à ceci près que tous les facteurs voisins de i sont pris en compte. Le facteur de normalisation $2^{-F_{i+a \in \partial i}}$ correspond à la contribution de la variable i , ainsi que ses facteurs voisins, à l'énergie libre totale (cf. figure 4.3a). Explicitement :

$$F_{i+a \in \partial i} = \hat{F}_{\circ + \square \in \circ}(\{q_{a \rightarrow i}\}_{a \in \partial i}) \doteq -\log \sum_{\sigma_i} \prod_{a \in \partial i} q_{a \rightarrow i}(\sigma_i), \quad (4.20)$$

où $\hat{F}_{\circ + \square \in \circ}$ est une fonction générique des messages $q_{a \rightarrow i}$ destinés à i , qui donne la contribution d'énergie libre d'une variable et de ses facteurs voisins. Comme dans le cas de la chaîne d'Ising, la somme sur i de ces contributions compte chaque facteur autant de fois qu'il a de voisins. Afin de corriger cela, il faut estimer la contribution

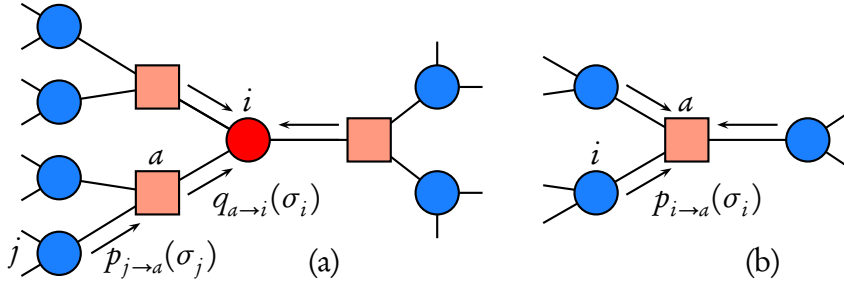


Fig. 4.3: Contributions à l'énergie libre. (a) Le branchement des sous-arbres de racine j par le rétablissement fictif des facteurs $a \in \partial i$, multiplie la fonction de partition par la constante de normalisation $2^{-F_{i+a \in \partial i}}$, cf. (4.20). (b) La contribution du facteur a s'obtient également comme l'effet multiplicatif que son rétablissement engendre sur la fonction de partition, cf. (4.21).

d'un seul facteur a en fonction des marginales de cavité. On considère le branchement, par le rétablissement fictif du facteur a , des sous-arbres attachés à ses voisins $i \in \partial a$ (cf. figure 4.3b). À l'issue de ce branchement, la fonction de partition est multipliée par 2^{-F_a} , avec :

$$F_a = \hat{F}_{\square}(\chi_a, \{p_{i \rightarrow a}\}_{a \in \partial i}) \doteq -\log \sum_{\sigma_a} \chi_a(\sigma_a) \prod_{i \in \partial a} p_{i \rightarrow a}(\sigma_i). \quad (4.21)$$

L'énergie libre totale vaut alors :

$$F = -\log Z = \sum_{i=1}^N F_{i+a \in \partial i} - \sum_{a=1}^M (|\partial a| - 1) F_a \quad (4.22)$$

Il est intéressant de remarquer que ces résultats exacts peuvent s'obtenir de manière équivalente par une approche markovienne relevant de la même logique qu'au paragraphe 1.1.4. Le processus de Markov consiste alors en une série de branchements successifs d'arbres : à chaque étape, plusieurs sous-arbres de racines $j \in J$ sont connectés par l'ajout d'une nouvelle racine i . La variable σ_i ne dépend que des $\{\sigma_j\}_{j \in J}$, et on suppose en outre que la loi conditionnelle de σ_i se factorise suivant une partition de $J = \bigcap_a \partial a$, de sorte que :

$$q_i(\sigma_i | \{\sigma_j\}_{j \in J}) = \prod_a \tilde{\chi}_a(\sigma_i, \{\sigma_j\}_{j \in \partial a}). \quad (4.23)$$

Le réglage des fonctions $\tilde{\chi}_a$ en vue d'obtenir une mesure globale de la forme (4.16) conduit précisément aux équations de cavité. Notez qu'il existe autant de manières d'opérer cette correspondance qu'il existe de choix pour la racine finale dans le processus de Markov. En cela, la méthode de la cavité est plus souple, car elle est indifférente à la notion de racine, et de sens de parcours.

4.1.3 Extension aux graphes dilués

Bien que la méthode de la cavité ne soit exacte que sur les arbres, elle peut s'avérer très efficace sur des hypergraphes aléatoires dilués, sous certaines conditions. En effet, de tels graphes ressemblent localement à des arbres quand le nombre de nœuds devient grand : partant d'une variable i donnée, et considérant les générations successives de ses voisins, la probabilité que i lui-même soit un voisin de génération au plus d vaut :

$$1 - \exp \left[\ln \left(1 - \frac{1}{N} \right) \sum_{g=1}^d [\rho'(1)\lambda'(1)]^g \right], \quad (4.24)$$

Cette probabilité reste presque sûrement nulle tant que :

$$d \ll \frac{\ln N}{\ln [\rho'(1)\lambda'(1)]}. \quad (4.25)$$

La taille typique des cycles de l'hypergraphe se comporte donc comme le logarithme de sa taille.

Dans la récursion de cavité sur les arbres, nous avons justifié que la loi jointe des seconds voisins d'une variable i pouvaient être factorisée *en l'absence* i . Dans un hypergraphe aléatoire, deux de ces seconds voisins seront typiquement séparés, si l'on fait abstraction du site i lui-même, par un chemin de longueur $\Theta(\ln N)$. Cette observation ne suffit cependant pas à assurer la validité de l'approximation des arbres. Il faut également veiller à ce que la dépendance entre deux variables quelconques décroisse suffisamment rapidement avec leur distance. Par exemple, une décroissance exponentielle des fonctions de corrélation :

$$\left| \mathbb{E}(\sigma_j \sigma_k) - \mathbb{E}\sigma_j \mathbb{E}\sigma_k \right| \leq e^{-\gamma d_{jk}}, \quad (4.26)$$

où d_{jk} est la longueur du plus court chemin reliant j et k , et ignorant i , suffit à valider l'hypothèse d'indépendance entre j et k , quand $N \rightarrow \infty$. Elle ne règle cependant pas la question de la validité de l'approximation de Bethe dans son ensemble, qui requiert en sus la prise en compte des effets d'interdépendance collective, et des propriétés de reconstruction. Nous discuterons plus en détail au paragraphe 4.1.6 les différents critères mettant à l'épreuve l'approximation des arbres.

Chaque solution des équations de Bethe décrit un *état thermodynamique* du système, c'est-à-dire une sous-partie c des configurations σ régies par la mesure induite $\propto p(\sigma)\mathbb{I}(\sigma \in c)$, sous laquelle les fonctions de corrélation décroissent rapidement. Quand cette solution est unique, on parle de phase liquide, tandis que l'existence d'un nombre exponentiel de solutions signale la présence d'une phase vitreuse. Le chapitre 6 présente une méthode générale permettant de traiter la méthode de la cavité avec multiplicité d'états.

L'émergence d'une multiplicité d'états est intimement lié à fragmentation dans les problèmes difficiles de satisfaction de contraintes comme k -SAT. En effet, à chaque composante connexe de l'espace des solutions, il est naturel d'associer une solution aux équations de Bethe. Les notions d'amas et d'états, bien que proches, ne sont cependant pas identiques, comme en témoigne la phase liquide fragmentée de k -SAT, cf. 2.3.2 : dans cette phase, à chaque amas est associé un état propre, mais il existe en plus un « super-état » englobant tous les amas et rendant compte de la mesure totale.

Variationnalité

Mentionnons une approche alternative à la méthode de la cavité, qui conduit également à l'approximation de Bethe. Dans cette approche [Yed01], la mesure de probabilité $p(\sigma)$ est approximée par une mesure-test factorisée, exacte sur les arbres :

$$p(\sigma) = \prod_a p_{\partial a}(\sigma_a) \prod_i p_i(\sigma_i)^{1-|\partial i|}, \quad (4.27)$$

où $p_{\partial a}$ est la loi jointe marginale des voisins de a . Cette mesure-test est insérée dans l'expression de l'énergie libre fonctionnelle de Gibbs :

$$G[p(\sigma)] = - \sum_{\sigma} p(\sigma) \log \prod_a \chi_a(\sigma_a) + \sum_{\sigma} p(\sigma) \log p(\sigma), \quad (4.28)$$

qui est minimisée avec des paramètres de Lagrange assurant la cohérence entre les marginales, ainsi que leur normalisation. La mise en œuvre de cette minimisation est, à un changement de variables près, équivalente à la méthode de la cavité.

Une des vertus de cette approche est qu'elle met en évidence le caractère *variationnel* des équations de cavité. Plus précisément, revenant maintenant au formalisme de la cavité, si l'on écrit l'énergie libre totale sous la forme :

$$F(\{p_{i \rightarrow a}, q_{a \rightarrow i}\}) = \sum_i F_{i+a \in \partial i}(\{q_{a \rightarrow i}\}_{a \in \partial i}) + \sum_a F_a(\{p_{i \rightarrow a}\}_{i \in \partial a}) - \sum_{(i,a)} F_{ai}(p_{i \rightarrow a}, q_{a \rightarrow i}), \quad (4.29)$$

avec $F_{ai} = -\log \sum_{\sigma_i} p_{i \rightarrow a}(\sigma_i) q_{a \rightarrow i}(\sigma_i)$, un rapide calcul montre cette fonctionnelle d'énergie libre est *stationnaire* dès que les équations de cavité (4.17), (4.18) sont vérifiées :

$$\forall (i,a), \quad \frac{\partial F}{\partial p_{i \rightarrow a}} = 0 \quad \frac{\partial F}{\partial q_{a \rightarrow i}} = 0. \quad (4.30)$$

Cette stationnarité a une conséquence pratique appréciable. Elle permet, quand la définition de la mesure $p(\sigma)$ dépend d'un paramètre externe λ (température, potentiel chimique, etc.), d'effectuer la dérivation de F par rapport à ce paramètre en ne tenant compte que de la dérivée explicite :

$$\frac{dF}{d\lambda} = \frac{\partial F}{\partial \lambda} + \sum_{(i,a)} \left[\frac{\partial F}{\partial p_{i \rightarrow a}} \frac{\partial p_{i \rightarrow a}}{\partial \lambda} + \frac{\partial F}{\partial q_{a \rightarrow i}} \frac{\partial q_{a \rightarrow i}}{\partial \lambda} \right] = \frac{\partial F}{\partial \lambda}. \quad (4.31)$$

Cela sera particulièrement utile dans toutes les situations faisant intervenir des transformations de Legendre.

4.1.4 Propagation des convictions

Afin de résoudre les équations de cavité, on peut simplement les implémenter en tant que formules itératives :

$$p_{i \rightarrow a}^{t+1}(\sigma_i) = \hat{p} \left[\{q_{b \rightarrow i}^t\}_{b \in i \setminus a} \right] (\sigma_i), \quad (4.32)$$

$$q_{b \rightarrow i}^t(\sigma_i) = \hat{q} \left[\chi_b, \{p_{j \rightarrow b}^t\}_{j \in b \setminus i} \right] (\sigma_i), \quad (4.33)$$

à partir de conditions initiales arbitraires. Cette itération définit l'algorithme de propagation des convictions (*Belief Propagation*, BP). Dans cet algorithme, les variables et les facteurs échangent des messages, $p_{i \rightarrow a}^t$ et $q_{a \rightarrow i}^t$, appelés *convictions*, qui contiennent une information locale sur le système :

- $p_{i \rightarrow a}^t$: la variable i renseigne le facteur a sur sa loi en l'absence de celui-ci.
- $q_{a \rightarrow i}^t$: le facteur a donne son avis sur la loi de σ_i .

Les itérations de l'algorithme BP permettent à chaque conviction d'être actualisée en fonction des nouvelles informations reçues. BP tend donc à « mettre d'accord » les agents en jeu en harmonisant les convictions, et à atteindre ainsi un consensus cohérent. Il peut aussi arriver que plusieurs consensus coexistent. Dans ce cas, soit l'algorithme converge vers l'un d'entre eux, soit il échoue.

L'algorithme BP décrit par les équations (4.32), (4.33) implémente une actualisation parallèle des messages. Il est possible de relaxer le processus itératif en actualisant les convictions une par une, dans un ordre séquentiel aléatoire. Une telle relaxation améliore en général la convergence de l'algorithme.

4.1.5 Statistique sur les instances

L'analyse statistique du comportement de la solution de cavité sur un ensemble d'instances permet théoriquement, si l'on admet la validité de l'approximation de Bethe dans la limite des grandes tailles, de calculer des quantités globales automoyennantes comme l'énergie libre, l'énergie moyenne, l'entropie, la magnétisation ou la distribution des recouvrements. Une instance aléatoire est définie par :

- un graphe factoriel dilué aléatoire, caractérisé par ses distributions de degrés $L(\ell)$ et $R(k)$.
- des facteurs χ_a , tirés au hasard selon une distribution dépendant du problème considéré. Par exemple, dans k -SAT, le facteur χ_a peut prendre uniformément 2^k valeurs possibles, suivant que ses variables sont niées ou pas.

Dans la limite des grandes tailles ($N \rightarrow \infty$), la densité de probabilité des marginales de cavité par rapport au choix aléatoire et uniforme d'une arête :

$$P(p) \doteq \frac{\mathbb{P}(p_{i \rightarrow a} \in [p, p + dp])}{dp}, \quad (4.34)$$

est donnée par le comportement à la racine d'un arbre aléatoire infini. Cette loi satisfait le système d'équations fermées :

$$P(p) = \sum_{\ell} \lambda(\ell) \int Q(q_1) dq_1 \cdots Q(q_{\ell}) dq_{\ell} \mathbb{E}_{\chi} \{ \delta [p - \hat{p}(\chi, q_1, \dots, q_{\ell})] \} \quad (4.35)$$

$$Q(q) = \sum_k \rho(k) \int P(p_1) dp_1 \cdots P(p_k) dp_k \delta [q - \hat{q}(p_1, \dots, p_k)] \quad (4.36)$$

Les fonctions $\hat{p}$ et $\hat{q}$ sont les fonctions universelles d'itérations de BP (4.17), (4.18). $Q(q)$ désigne la densité de probabilité des messages $q_{a \rightarrow i}$. Rappelons que les distributions $\lambda(\ell)$ et $\rho(k)$ correspondent aux distributions de degrés d'un lien (i, a) tiré au hasard (cf. §3.1.3). La distribution $P(p)$ des marginales de cavité étant connues, on peut en déduire la distribution des marginales totales :

$$\frac{\mathbb{P}(p_i \in [p, p + dp])}{dp} = \sum_{\ell} L(\ell) \int Q(q_1) dq_1 \cdots Q(q_{\ell}) dq_{\ell} \mathbb{E}_{\chi} \{ \delta [p - \hat{p}(\chi, q_1, \dots, q_{\ell})] \}, \quad (4.37)$$

ainsi que l'énergie libre réduite :

$$\begin{aligned} f = \lim_{N \rightarrow \infty} \frac{F}{N} &= \sum_{\ell} L(\ell) \int Q(q_1) dq_1 \cdots Q(q_{\ell}) dq_{\ell} \hat{F}_{\circ + \square \in \circ}(q_1, \dots, q_{\ell}) \\ &\quad - \frac{\mathbb{E}(\ell)}{\mathbb{E}(k)} \sum_k R(k)(k-1) \int P(p_1) dp_1 \cdots P(p_k) dp_k \mathbb{E}_{\chi} [\hat{F}_{\square}(\chi, p_1, \dots, p_k)], \end{aligned} \quad (4.38)$$

où les fonctions $\hat{F}_{\circ + \square \in \circ}$ et $\hat{F}_{\square}$ sont définies équations (4.20), (4.21). Pour les mêmes raisons que dans le cas d'une instance donnée, f est stationnaire en tant que fonction de $P(p)$.

Ces équations moyennées sont équivalentes aux équations de col obtenues par la méthode des répliques sous l'Ansatz de symétrie des répliques [Mon98]. Elle peuvent être résolues par la « dynamique des populations » [MP01]. Dans la mise en œuvre de cette technique, les distributions $P(p)$ et $Q(q)$ sont représentées par deux grandes collections de nombres, appelées populations, que l'on actualise par le renouvellement progressif de leurs individus selon $\hat{p}$ et $\hat{q}$.

4.1.6 Stabilité et restructurabilité

De quels critères dispose-t-on pour vérifier la validité de l'approximation de Bethe ? Plusieurs approches ont été proposées à cette fin, dont certaines s'avèrent équivalentes. Nous faisons rarement recours à ces critères dans cette thèse. Aussi nous contenterons-nous de les exposer brièvement.

Une première approche consiste à envisager la méthode de la cavité à un état unique comme un cas limite d'un cadre plus général, où est postulée la multiplicité d'états. Dans ce contexte, la cohérence interne de la solution à un état unique peut être éprouvée par l'étude de sa stabilité au sein d'un espace à états multiples. Cette voie, suivie par [MRT03, MPRT04], permet de délimiter la zone de stabilité de la solution à un état unique. En étendant le raisonnement à un niveau hiérarchique supplémentaire, la stabilité des équations avec multiplicité « simple » d'états peut être testée dans le cadre plus général d'une multiplicité « double », où les états eux-mêmes sont regroupés dans des super-états¹. Il faut néanmoins noter que si ce critère garantit la cohérence interne de la solution, il n'exclut pas qu'une autre solution thermodynamiquement plus favorable la supplante à l'issue d'une transition du premier ordre.

La stabilité des solutions peut également être testée localement, sans recourir à un espace plus grand. En vertu du principe de fluctuation-dissipation, on peut montrer que l'instabilité de l'état unique est équivalente à la divergence de la susceptibilité de verre :

$$\chi_2 = \frac{1}{N} \sum_{i,j} \left[\mathbb{E}(\sigma_i \sigma_j) - \mathbb{E}(\sigma_i) \mathbb{E}(\sigma_j) \right]^2 \quad (4.39)$$

Cette équivalence est établie dans [RBMM04], pour les solutions à un état unique comme pour les solutions avec un nombre arbitraire de niveaux hiérarchiques d'états. À titre d'exemple, ces méthodes ont permis de calculer la fenêtre de stabilité de la solution de cavité avec multiplicité d'états dans le problème k -SAT [MMZ06]. Le fait que cette fenêtre contient le seuil de satisfaisabilité $\alpha_s(k)$ soutient la validité de la prédiction de ce seuil par la méthode de la cavité (voir §6.1.4).

Une approche alternative réside dans l'étude des propriétés de reconstruction des arbres contenus dans le graphe [MM06a]. Connaissant la valeur d'une variable à la racine d'un arbre, quelle information contiennent les branches quand le nombre de générations devient grand ? L'information se dissipe-t-elle rapidement dans le graphe ? L'hypothèse de l'état unique suppose en effet une propriété d'« amnésie » qui assure l'indépendance vis-à-vis des conditions aux bords. La reformulation de cette propriété en termes de corrélations requiert la notion de *fonction de corrélation entre un point et un ensemble* [MS06b] : cette fonction quantifie la capacité de relaxation d'une variable, consécutivement au gel de ses voisins de $g^{\text{ième}}$ génération. Comme auparavant,

¹Ce sont les schémas de brisure de symétrie des répliques à un et deux pas, respectivement.

ce critère s'avère équivalent à la condition de stabilité.

Enfin, on peut envisager de corriger la méthode de la cavité en prenant explicitement en compte les boucles du graphe au delà de l'approximation des arbres, par un jeu d'approximation de plus en plus précises. C'est l'approche adoptée par [MR05] et par [CC06]. Bien qu'intuitivement plus naturelle, cette approche s'avère moins fructueuse, car elle ne permet pas de rendre compte des effets collectifs propres aux phases vitreuses.

4.2 Exemples

4.2.1 Décodage itératif

Voyons maintenant un exemple simple et utile d'application de la méthode de la cavité : l'étude de la performance d'un code LDPC sur le canal d'effacement. Supposons qu'un mot de code σ^0 soit transmis par le canal d'effacement, et notons E l'ensemble des bits effacés. Le code étant linéaire, on peut se ramener à $\sigma^0 = 0$ sans perte de généralité. Dans la phase de décodage, la mesure *a posteriori* s'écrit :

$$p(\sigma) = \frac{1}{Z} \prod_{a=1}^M \mathbb{I} \left(\sum_{i \in \partial a \cap E} \sigma_i = 0 \right) \prod_{i \notin E} \mathbb{I}(\sigma_i = 0) \quad (4.40)$$

Si $Z > 1$, le décodage échoue ; sinon, il réussit.

Les équations de propagation des convictions s'écrivent :

$$p_{i \rightarrow a}^{t+1}(\sigma_i) \propto \begin{cases} \prod_{b \in i \setminus a} q_{b \rightarrow i}^t(\sigma_i) & \text{si } i \in E, \\ \delta_{\sigma_i, 0} & \text{sinon.} \end{cases} \quad (4.41)$$

$$q_{b \rightarrow i}^t(\sigma_i) = \frac{1}{2} \left[1 + \sigma_i \prod_{j \in \partial b \setminus i} \left(p_{j \rightarrow b}^t(0) - p_{j \rightarrow b}^t(1) \right) \right] \quad (4.42)$$

Si l'on prend pour conditions initiales $p_{i \rightarrow a}^0 = \frac{1}{2}(\delta_0 + \delta_1) \doteq \gamma$, il est facile de vérifier que tous les messages valent à tout temps γ ou δ_0 . Cette observation entraîne la simplification des règles de cavité :

$$p_{i \rightarrow a}^{t+1} = \begin{cases} \delta_0 & \text{si } \exists b \in \partial i \setminus a \text{ tel que } q_{b \rightarrow i}^t = \delta_0 \text{ ou si } i \notin E, \\ \gamma & \text{sinon.} \end{cases} \quad (4.43)$$

$$q_{b \rightarrow i}^t = \begin{cases} \delta_0 & \text{si } \forall j \in \partial b \setminus i, p_{j \rightarrow b}^t = \delta_0 \\ \gamma & \text{sinon} \end{cases} \quad (4.44)$$

Il se trouve que cet algorithme BP procède aux mêmes opérations que l'algorithme d'effeuillage décrit au paragraphe 3.2.3. La transmission du message $q_{b \rightarrow i} = \delta_0$ correspond à l'opération d'effeuillage proprement dite, par laquelle l'incertitude sur la variable i est levée grâce au test de parité b , cf. (4.43). Un tel message ne sera envoyé que si tous les autres voisins de b sont eux-mêmes connus, cf. (4.44).

Notons $\{p_{i \rightarrow a}, q_{a \rightarrow i}\}$ le point fixe des équations BP. La formule générale de Bethe pour l'énergie libre (4.22) permet ici d'estimer l'entropie :

$$S = \log Z = \sum_{i \in E} \prod_{a \in \partial i} \delta(q_{a \rightarrow i}, \gamma) - \sum_{(i,a)} \delta(q_{a \rightarrow i}, \gamma) + \sum_{a=1}^M (k_a - 1) \left[1 - \prod_{i \in \partial a} \delta(p_{i \rightarrow a}, \delta_0) \right] \quad (4.45)$$

La statistique des messages sur un code aléatoire s'effectue de manière analytique. Soit un ensemble de codes aléatoires caractérisé par ses distributions de degrés (L, R) . On note η^t la probabilité qu'un message $p_{i \rightarrow a}^t$ choisi au hasard vaille γ , et ζ^t la probabilité que $q_{a \rightarrow i}^t = \gamma$. Dans la limite des grands mots, la transposition de (4.35) et (4.36) donne :

$$\eta^{t+1} = \epsilon \sum_{\ell} \lambda(\ell) (\zeta^t)^\ell = \epsilon \lambda(\zeta^t), \quad (4.46)$$

$$\zeta^t = 1 - \sum_k \rho(k) (1 - \eta^t)^k = 1 - \rho(1 - \eta^t). \quad (4.47)$$

Au point fixe (η, ζ) de ces équations, l'entropie réduite, cf. (4.38), s'évalue à

$$s = \lim_{N \rightarrow \infty} \frac{S}{N} = \epsilon L(\zeta) - \frac{L'(1)}{R'(1)} [1 - R(1 - \eta) - \eta R'(1 - \eta)]. \quad (4.48)$$

Trois comportements se dégagent suivant la valeur du bruit (voir figure 4.4). Quand $\epsilon < \epsilon_d$, la seule solution aux équations de cavité est triviale $(\eta, \zeta) = (0, 0)$: l'algorithme BP a su recouvrer le mot entier. Quand $\epsilon_d < \epsilon < \epsilon_c$, l'algorithme BP est bloqué par un point fixe non-trivial, causé par la présence d'une sous-partie d'arrêt. Pour autant, l'évaluation de l'entropie (4.48) donne une valeur négative. Ceci est surprenant, car on sait que l'entropie est toujours positive ou nulle. Mais d'une part, ce point fixe « ignore »² le mot de code original $(0, \dots, 0)$, qui fait l'objet d'une solution distincte des équations de cavité. D'autre part, on sait qu'une entropie négative signale en fait un Z typiquement nul, et s'explique par des événements rares [Riv04]. Le bruit critique ϵ_c pour lequel l'entropie devient strictement positive définit le seuil de décodabilité optimal.

²En termes physiques, l'algorithme vu comme processus d'évolution local reste bloqué dans une phase vitreuse et ignore le cristal.

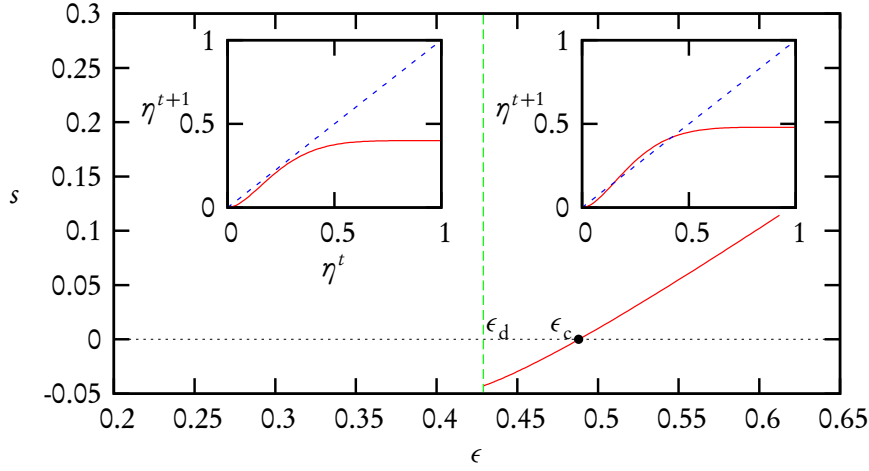


Fig. 4.4: Diagramme de phases du code régulier ($k = 6, \ell = 3$) sur le canal d'effacement. L'entropie (4.48) est représentée en fonction de ϵ . Dans les encarts sont représentées les équations d'évolution de BP, pour $\epsilon = 0,4$ (gauche) et $\epsilon = 0,48$ (droite). Quand $\epsilon < \epsilon_d \approx 0,42944$, l'unique point fixe est trivial, traduisant le fait que BP résout le problème entier. L'entropie obtenue par la cavité devient positive pour $\epsilon > \epsilon_c \approx 0,48815$.

Le seuil itératif ϵ_d a originellement été dérivé dans [LMS⁺97], et le seuil optimal a été obtenu pour la première fois par la méthode des répliques [FLMRT02]. Ces prédictions ont depuis fait l'objet d'une preuve rigoureuse [MMRU04] basée sur l'analyse de l'algorithme BP et la validation de l'approximation des arbres.

L'analyse permet aussi de mettre en évidence une observation importante, également valable pour le BSC : à taux fixé, les codes affichant les meilleures performances itératives sont irréguliers. Pour preuve, l'ensemble de taux $R = 1/2$ défini par [RU07] :

$$\begin{aligned}\lambda(x) &= 0,106257x + 0,486659x^2 + 0,010390x^{10} + 0,396694x^{19}, \\ \rho(x) &= 0,5x^7 + 0,5x^8,\end{aligned}\tag{4.49}$$

bénéficie d'un seuil critique itératif $\epsilon_d \approx 0,4741$ proche de la borne de Shannon $\epsilon_{\text{sh}} = 1 - R$, et bat tous les codes régulier de même taux. En fait, il est même possible de saturer la borne de Shannon avec des séquences de codes bien choisis [LMSS01].

En revanche, la limite des grandes connectivités n'offre pas une solution viable. Le seuil optimal ϵ_c tend bien vers la borne de Shannon, de manière cohérente avec l'étude du chapitre précédent, mais le seuil itératif ϵ_d tend quant à lui vers 0.

4.2.2 Énumération des A -parties d'un graphe factoriel

Nous illustrons maintenant l'approximation de Bethe par un exemple original faisant intervenir une énergie et une température : le décompte des A -parties, c'est-à-dire des solutions du problème d'occupation défini au paragraphe 3.3. Pour une partie A de $\mathbb{N}$,

$$S \text{ est une } A\text{-partie ssi pour chaque facteur } a, \quad |\partial a \cap S| \in A. \quad (4.50)$$

On veut calculer le nombre n_w de A -parties S de taille $|S| = w$. Ce calcul est très difficile en général : il est démontré que le problème de décision $n_w \stackrel{?}{=} 0$ est NP -complet dans les cas particuliers des mots de codes [Var97], des sous-parties d'arrêt [KS05] et de la couverture de graphes [GJ79], p.190.

Avant de pouvoir utiliser l'approximation des arbres, il convient de reformuler le problème dans l'ensemble canonique, en définissant la mesure :

$$p(\sigma, \beta) = \frac{1}{Z(\beta)} \prod_{a=1}^M \mathbb{I} \left(\sum_{i \in \partial a} |\sigma_i| \in A \right) 2^{-\beta \sum_{i=1}^N |\sigma_i|}, \quad (4.51)$$

où $|S| = \sum_i |\sigma_i|$ joue ici le rôle d'une énergie. La fonction de partition s'identifie à la fonction génératrice de la séquence n_w :

$$Z(\beta) = \sum_w n_w 2^{-\beta w}, \quad (4.52)$$

appelée *fonction d'énumération des A -parties*, et qui généralise la fonction d'énumération des poids ($A = 2\mathbb{N}$, cf. (1.60)). On est ainsi ramené à un problème de physique statistique, que l'on peut traiter par la méthode de la cavité. Les deux quantités $W(w) = \log n_w$, et $\Phi(\beta) = -\log Z(\beta)$ se déduisent l'une de l'autre, dans la limite des grands N , par des transformations de Legendre :

$$W(w) \approx \beta w - \Phi(\beta), \quad \text{avec } w = \partial_\beta \Phi(\beta). \quad (4.53)$$

La formule de Bethe (4.22) permet d'évaluer $\Phi(\beta)$, d'où l'on tire $W(w)$ sous forme paramétrée :

$$\begin{aligned} W(w) \approx & \beta w + \sum_{i=1}^N \log \left(1 + 2^{-\beta} \prod_{a \in i} y_{a \rightarrow i} \right) \\ & + \sum_{a=1}^M \log \left(\sum_{J \subset \partial a, |J| \in A} \prod_{i \in J} x_{i \rightarrow a} \right) - \sum_{(i,a)} \log(1 + x_{i \rightarrow a} y_{a \rightarrow i}) \end{aligned} \quad (4.54)$$

$$w = \sum_{i=1}^N \frac{2^{-\beta} \prod_{a \in i} y_{a \rightarrow i}}{1 + 2^{-\beta} \prod_{a \in i} y_{a \rightarrow i}} \quad (4.55)$$

où $x_{i \rightarrow a} \doteq p_{i \rightarrow a}(1)/p_{i \rightarrow a}(0)$ et $y_{a \rightarrow i} \doteq q_{a \rightarrow i}(1)/p_{a \rightarrow i}(0)$ vérifient les équations de point fixe :

$$x_{i \rightarrow a} = 2^{-\beta} \prod_{b \in i \setminus a} y_{b \rightarrow i} \quad (4.56)$$

$$y_{b \rightarrow i} = \frac{\sum_{J \subset \partial b \setminus i} \prod_{j \in J} x_{j \rightarrow b}}{\sum_{J \subset \partial b \setminus i} \prod_{j \in J} x_{j \rightarrow b}}. \quad (4.57)$$

La quantité $W(w)/N$ est, comme à l'accoutumée, supposée automoyennante. Par souci de concision, nous omettons de reproduire ici les équations de cavité moyennées (4.35)–(4.38) permettant d'en évaluer la limite.

Dans le cas d'un graphe factoriel régulier (ℓ, k) , le graphe est localement doté d'une structure d'arbre régulier. Cette invariance translationnelle permet de proposer un Ansatz factorisé, où les messages sont constants : $x_{i \rightarrow a} = x$, $y_{a \rightarrow i} = y$. On peut alors estimer analytiquement l'entropie microcanonique, aussi appelée « taux de croissance » de n_w :

$$\Omega(\omega) = \lim_{N \rightarrow \infty} \frac{W(N\omega)}{N} = (1 - \ell)H(\omega) - \omega \ell \log(x) + \frac{\ell}{k} \log \left[\sum_{d \in A} \binom{k}{d} x^d \right] \quad (4.58)$$

$$\omega = x \frac{\sum_{d+1 \in A} \binom{k-1}{d} x^d}{\sum_{d=0}^k \binom{k}{d} x^d}. \quad (4.59)$$

Nous illustrons ce calcul sur l'exemple des sous-parties d'arrêt. La figure 4.5 présente l'entropie microcanonique Ω en fonction de la taille réduite $\omega = w/N$, pour deux constructions de codes. Ici encore, les régions d'entropie négatives ne sont pas physiques : dans ces régions le nombre de sous-parties d'arrêt est typiquement nul. La taille minimale $\omega_{\min} > 0$ de sous-partie d'arrêt non-triviale s'obtient comme la plus petite racine de $\Omega(\omega)$, quand celle-ci existe. Les petites sous-parties d'arrêt sont la cause la plus fréquente d'échec du décodage itératif quand le niveau de bruit est faible. Elles expliquent [KV03, OVZ05] le fameux « plancher d'erreur » (*error floor*) observé dans les codes linéaires [MP03], cf. figure 4.6.

Le panneau de droite de la figure 4.5 représente l'exposant Ω pour le code irrégulier de taux $R = 1/2$ défini par l'équation (4.49). Bien que cet ensemble soit très performant du point de vue algorithmique, on n'y observe pas de « fossé » dans les tailles possibles de sous-parties d'arrêt. Cela n'a en fait rien de contradictoire avec les remarques précédentes : si ce code affiche de bonnes performances près du bruit critique ϵ_d dans la région dite de « cascade » (*waterfall*), son plancher d'erreur (*error floor*) est en revanche assez haut. Parmi les autres caractéristiques de cet ensemble, on observe une transition du premier ordre entre deux solutions de cavité. Dans le régime où ces deux solutions coexistent, celle d'entropie maximale domine.

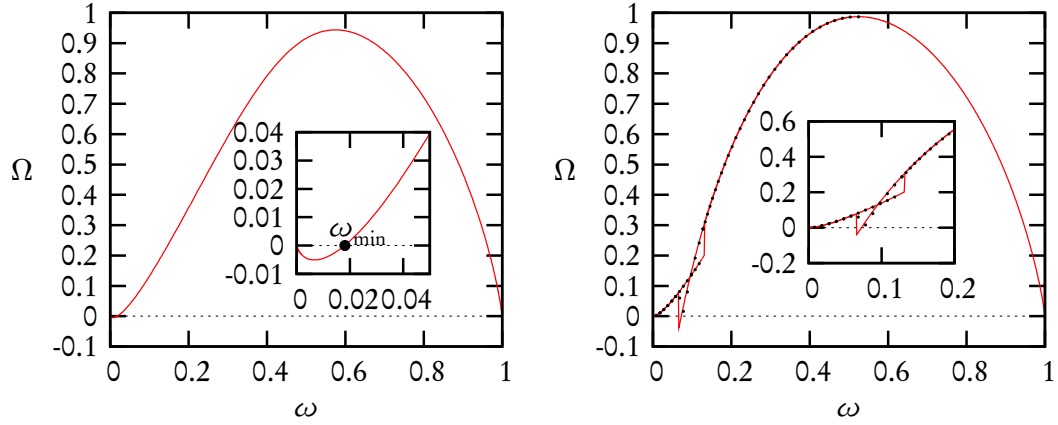


Fig. 4.5: Entropie gouvernant le nombre $2^{N\Omega}$ de sous-parties d'arrêt en fonction de leur taille $|S| = \omega N$, pour deux constructions de code. À gauche, l'ensemble régulier ($k = 6, \ell = 3$) : on y observe une région de petites tailles interdites ($\omega < \omega_{\min}$). Dans l'ensemble optimisé, cf. (4.49), à droite, cette région n'existe pas. En revanche l'exposant subit une transition de phase du premier ordre. Sur le graphe de droite, les points correspondent aux équations de cavité moyennées, tandis que la ligne pleine — représente le résultat de la moyenne *recuite*, dont le calcul est détaillé au paragraphe 5.1. L'encart à droite donne un plan rapproché sur la région des petits ω , permettant de distinguer ces deux moyennes. Dans le cas régulier en revanche (à gauche) les deux moyennes coïncident.

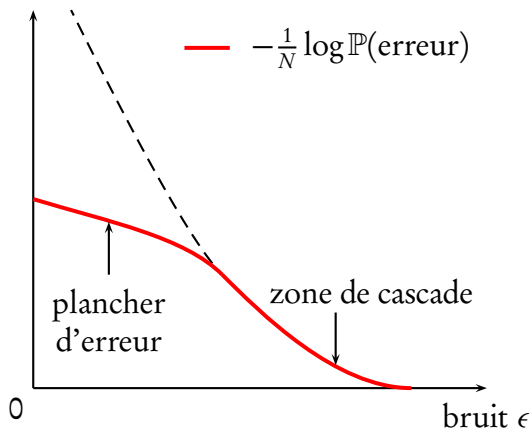


Fig. 4.6: Le plancher d'erreur dans les codes. Près du seuil de décodabilité, l'exposant d'erreur croît rapidement, dans une région dite de cascade. Pour les plus niveaux plus faibles de bruit en revanche, l'exposant est limité par un plancher d'erreur causé par les sous-parties d'arrêt de faible poids (les termes de plancher et de cascade prennent leur sens si l'on considère, comme les informaticiens le font, l'opposé de l'exposant d'erreur).

4.3 Calcul des corrélations

4.3.1 Propagation des susceptibilités

La méthode de cavité, dont l'une des fonctions est d'approximer les marginales dans les modèles graphiques, peut facilement être étendue aux calculs des fonctions de corrélation. Celles-ci peuvent se déduire des fonctions de réponse à l'aide de la relation de fluctuation-dissipation. Afin d'évaluer ces réponses, on applique sur chaque variable un champ extérieur h_i^σ :

$$p(\sigma, \{h_i\}) = \frac{1}{Z(\{h_i\})} \prod_{a=1}^M \chi_a(\sigma_a) \prod_{i=1}^N \prod_{\sigma} 2^{h_i^\sigma \delta_{\sigma_i, \sigma}}. \quad (4.60)$$

Le champ extérieur est en fait un vecteur à q éléments, où q est la taille de l'alphabet de σ_i . La relation fluctuation-dissipation établit la correspondance entre la fonction de corrélation connexe à deux points et la susceptibilité à champ nul :

$$\left. \frac{\partial p_i(\sigma_i)}{\partial h_j^{\sigma_j}} \right|_{h=0} = p_{ij}(\sigma_i, \sigma_j) - p_i(\sigma_i)p_j(\sigma_j) \doteq \pi_{ij}(\sigma_i, \sigma_j). \quad (4.61)$$

On définit des *susceptibilités de cavité* :

$$\pi_{i \rightarrow a, j}(\sigma_i, \sigma_j) = \left. \frac{\partial p_{i \rightarrow a}}{\partial h_j^{\sigma_j}} \right|_{h=0}, \quad \tilde{\pi}_{a \rightarrow i, j}(\sigma_i, \sigma_j) = \left. \frac{\partial q_{a \rightarrow i}}{\partial h_j^{\sigma_j}} \right|_{h=0}, \quad (4.62)$$

que l'on actualise à l'aide d'équations de *propagation des susceptibilités*, qui ne sont rien d'autre que les dérivées des équations de cavité :

$$\pi_{i \rightarrow a, j}(\sigma_i, \sigma_j) = p_{i \rightarrow a}(\sigma_i) \left[c_{i \rightarrow a, j}(\sigma_j) + \sum_{b \in \partial i \setminus a} \frac{\tilde{\pi}_{b \rightarrow i, j}(\sigma_i, \sigma_j)}{q_{b \rightarrow i}(\sigma_i)} + \delta_{i, j} \delta_{\sigma_i, \sigma_j} \right] \quad (4.63)$$

$$\tilde{\pi}_{b \rightarrow i, j}(\sigma_i, \sigma_j) = \sum_{k \in \partial b \setminus i} \sum_{\sigma_k} \frac{\partial \hat{q}(\sigma_i)}{\partial p_{k \rightarrow b}(\sigma_k)} \pi_{k \rightarrow b}(\sigma_k, \sigma_j), \quad (4.64)$$

où la fonction $\hat{q}$ est définie équation (4.18). La constante $c_{i \rightarrow a, j}(\sigma_j)$ est déterminée par la condition de normalisation :

$$\sum_{\sigma_i} \pi_{i \rightarrow a, j}(\sigma_i, \sigma_j) = 0. \quad (4.65)$$

La susceptibilité π_{ij} s'obtient alors comme :

$$\pi_{ij}(\sigma_i, \sigma_j) = p_i(\sigma_i) \left[c_{ij}(\sigma_j) + \sum_{a \in \partial i} \frac{\tilde{\pi}_{a \rightarrow i, j}(\sigma_i, \sigma_j)}{q_{a \rightarrow i}(\sigma_i)} + \delta_{i, j} \delta_{\sigma_i, \sigma_j} \right]. \quad (4.66)$$

Cette expression devient symétrique en i et en j dès lors que les équations de cavité sont vérifiées.

La propagation des susceptibilités, en ce qu'elle s'appuie sur l'approximation de Bethe, donne des résultats exacts sur les arbres, et demeure performante sur les grands graphes dilués. Le cas des arbres permet en outre quelques simplifications. La réponse de cavité par rapport à une perturbation en j ne peut affecter que les messages pointant dans la direction opposée à celle de j sur l'arbre. En effet, la définition des messages $i \rightarrow a$ pointant dans la direction de j suppose la suppression de la clause a , celle-là même qui relie i à j dans l'arbre. Considérons deux variables i_0 et i_n , reliées par le chemin (unique) sur l'arbre ($i_0 \rightarrow a_1 \rightarrow i_1 \rightarrow \dots \rightarrow a_n \rightarrow i_n$). Le message $\pi_{i_0 \rightarrow a_1, i_0}$ s'écrit, à une constance près, comme une fonction de Dirac. Le message suivant, $\tilde{\pi}_{a_0 \rightarrow i_1, i_0}$ ne dépend en fait que de $\pi_{i_0 \rightarrow a_0, i_0}$, car les autres messages pointant vers a_0 se propagent en direction de i_0 , et sont donc nuls. Le même argument peut être répété tout le long de la chaîne reliant i_0 à i_n , chaque nouveau message se déduisant uniquement du précédent sur la chaîne.

Dans un graphe plus général en revanche, même s'il est dilué, la propagation d'une réponse peut prendre une infinité de chemins différents, emprunter des boucles, etc. Le calcul de la susceptibilité impose donc de prendre en compte l'intégralité des messages.

La méthode présentée ici s'étend facilement aux fonctions de corrélation à n points, en recourant à des relations de fluctuation-dissipation généralisées. Les procédures algorithmiques associées, qui font intervenir des susceptibilités de cavité à $n + 1$ indices, deviennent de plus coûteuses à mesure que l'ordre des corrélations augmente.

4.3.2 Application : modèles d'entropie maximale

La propagation des susceptibilités peut servir à résoudre une large classe de problèmes d'apprentissage, basés sur le principe d'entropie maximale, et assimilables à des machines de Boltzmann [AHS87].

Dans certaines expériences de biologie, les données prennent la forme d'une série de valeurs discrètes σ à valeur dans un espace de grande dimension N , comme l'enregistrement simultané de l'activité d'un ensemble de neurones [SBSB06], la composition de protéines, la concentration jointe d'un ensemble de gènes, etc.

Dans certains cas, il est raisonnable de supposer que les éléments de cette série sont tirés indépendamment à partir d'une distribution inconnue $p(\sigma)$. L'échantillonnage complet de la mesure p requiert en principe un nombre exponentiel d'essais, ce qui devient assez vite impraticable quand N est grand. Cependant, on peut espérer comprendre la structure corrélative de p en l'approchant par une série de mesures-test, $p_1, p_2, \dots, p_N = p$, construites comme suit : p_n est la mesure d'entropie maximale

dont toutes les fonctions de corrélations d'ordre inférieur ou égal à n égalent exactement celles de p . La technique des multiplicateurs de Lagrange donne la forme que doit prendre p_n :

$$p_n(\sigma) = \frac{1}{Z_n} 2^{-E_n(\sigma)}, \quad \text{avec} \quad E_n(\sigma) = - \sum_{d=1}^n \sum_{i_1 < i_2 < \dots < i_d} J_{i_1 i_2 \dots i_d} \sigma_{i_1} \sigma_{i_2} \dots \sigma_{i_d}, \quad (4.67)$$

où l'on s'est restreint au cas de variables de spin $\sigma_i = \pm 1$. Cette mesure n'est autre qu'un modèle d'Ising généralisé, avec des interactions impliquant au plus n spins. Aux deux extrêmes de cette série d'approximations, p_1 correspond à un système de spins indépendants, et ne contient aucune information sur la structure corrélative des variables. À l'opposé, p_N est exactement égale à p . Afin de quantifier la précision d'une approximation intermédiaire p_n , on définit $I_n = H(p_1) - H(p_n)$ comme la « perte » d'entropie causée par la prise en compte des corrélations d'ordre au plus n au delà du modèle indépendant. Le rapport I_n/I_N donne la part de corrélations dont les fonctions à n points permettent de rendre compte. Plus ce rapport est proche de 1, meilleure est l'approximation.

Dans la suite nous nous restreignons au cas le plus simple, $n = 2$, ce qui revient à chercher un modèle d'Ising :

$$E(\sigma) = - \sum_{i < j} J_{ij} \sigma_i \sigma_j - \sum_i h_i \sigma_i \quad (4.68)$$

compatible avec les magnétisations locales et les fonctions de corrélations de paires du modèle original. On parle alors de problème d'Ising inverse ou de machine de Boltzmann. Le problème *direct*, qui consiste à déduire les fonctions de corrélation des couplages, est déjà un défi majeur du point de vue algorithmique : il ne peut généralement être résolu qu'en temps exponentiel ou, de manière approximative et sous l'hypothèse d'ergodicité, par des procédures de type Monte-Carlo. Il n'est donc pas surprenant que le problème inverse relève de la même difficulté.

La formulation originale du problème de la machine de Boltzmann [AHS87] prend en quelque sorte le chemin inverse de la définition que nous venons de donner, tout en y étant équivalente : étant donnée une mesure p , on cherche une mesure p_n de la forme (4.67) et minimisant la divergence de Kullback-Leibler avec p :

$$D(p||p_n) = \sum_{\sigma} p(\sigma) \log \frac{p(\sigma)}{p_n(\sigma)}. \quad (4.69)$$

Cette minimisation entraîne justement l'égalité des fonctions de corrélations jusqu'à l'ordre n . Une méthode standard de descente de gradient induit la règle d'apprentissage suivante (e. g. pour $n = 2$) :

$$J_{ij} \leftarrow J_{ij} - \epsilon \frac{\partial D(p||p_2)}{\partial J_{ij}} = J_{ij} + \epsilon (\chi_{ij} - \chi'_{ij}), \quad (4.70)$$

$$\text{où } \chi_{ij} = \mathbb{E}_p(\sigma_i \sigma_j) - \mathbb{E}_p(\sigma_i) \mathbb{E}_p(\sigma_j), \quad \chi'_{ij} = \mathbb{E}_{p_2}(\sigma_i \sigma_j) - \mathbb{E}_{p_2}(\sigma_i) \mathbb{E}_{p_2}(\sigma_j). \quad (4.71)$$

Cette règle d'apprentissage se heurte une fois de plus à la difficulté que représente l'estimation des susceptibilités χ'_{ij} dans le modèle d'Ising. Cette difficulté peut être partiellement contournée par l'utilisation d'algorithmes Monte-Carlo, par du recuit simulé, ou par des approximations de champ moyen [PA87, KR98, Tan98]. Notre but est ici de proposer une méthode alternative reposant sur l'approximation de Bethe, et plus précisément sur l'algorithme de propagation des susceptibilités.

Avant d'adapter les équations (4.63), (4.64) au modèle d'Ising (4.68), il convient de passer en convention de champs (appelés log-vraisemblances en théorie de l'information). On définit :

$$h_{i \rightarrow j} = \frac{1}{2} \log \frac{p_{i \rightarrow j}(+1)}{p_{i \rightarrow j}(-1)}, \quad u_{i \rightarrow j} = \frac{1}{2} \log \frac{q_{i \rightarrow j}(+1)}{q_{i \rightarrow j}(-1)}, \quad (4.72)$$

$$g_{i \rightarrow j, k} = \frac{\partial h_{i \rightarrow j}}{\partial h_k}, \quad v_{i \rightarrow j, k} = \frac{\partial u_{i \rightarrow j}}{\partial h_k} \quad (4.73)$$

La notation $i \rightarrow j$ est ici un raccourci pour $i \rightarrow a$, où a correspond au facteur J_{ij} . Les équations de cavité et leur dérivées prennent la forme :

$$h_{i \rightarrow j} = \sum_{k \in \partial i \setminus j} u_{k \rightarrow i} + h_i, \quad \tanh u_{k \rightarrow i} = \tanh J_{ik} \tanh h_{k \rightarrow i}, \quad (4.74)$$

$$g_{i \rightarrow j, k} = \sum_{l \in \partial i \setminus j} v_{l \rightarrow i, k} + \delta_{i, k}, \quad v_{l \rightarrow i, k} = g_{l \rightarrow i, k} \tanh J_{il} \frac{1 - \tanh^2 h_{l \rightarrow i}}{1 - \tanh^2 u_{l \rightarrow i}}. \quad (4.75)$$

Quand ces équations sont vérifiées, les fonctions de corrélation à un et deux points s'écrivent :

$$H_i \doteq \frac{1}{2} \log \frac{p_i(+1)}{p_i(-1)} = \sum_{j \in \partial i} u_{j \rightarrow i} + h_i \quad (4.76)$$

$$\chi_{ij} = \bar{\chi}_{ij} g_{j \rightarrow i, j} + g_{i \rightarrow j, j} (1 - \tanh^2 H_i), \quad (4.77)$$

où

$$\bar{\chi}_{ij} \doteq \frac{\tanh J_{ij} + \tanh h_{i \rightarrow j} \tanh h_{j \rightarrow i}}{1 + \tanh J_{ij} \tanh h_{i \rightarrow j} \tanh h_{j \rightarrow i}} - \tanh H_i \tanh H_j \quad (4.78)$$

On résout ainsi le problème « direct », en déduisant les fonctions de corrélation connexes et les champs effectifs locaux à partir des couplages J_{ij} et des champs extérieurs h_i . Dans le cas particulier de l'arbre, où ces équations sont exactes, la fonction de corrélation entre deux variables i et j reliées par le chemin $(i = i_0, i_1, \dots, i_n = j)$ se factorise :

$$\chi_{ij} = \frac{\prod_{a=1}^n \bar{\chi}_{i_{a-1} i_a}}{\prod_{a=1}^{n-1} (1 - \tanh^2 H_{i_a})}, \quad (4.79)$$

Il est relativement aisé, dans le cadre des équations d'échange de messages, de procéder à l'inversion du problème, en vue d'obtenir les couplages J_{ij} et les champs extérieurs h_i en fonction des fonctions de corrélations connexe à deux points χ_{ij} et des champs effectifs H_i . Il suffit pour cela de partir de l'équation (4.76), et de construire la règle d'actualisation :

$$\tanh J_{ij} \leftarrow \frac{\tilde{C}_{ij} - \tanh h_{i \rightarrow j} \tanh h_{j \rightarrow i}}{1 - \tilde{C}_{ij} \tanh h_{i \rightarrow j} \tanh h_{j \rightarrow i}}, \quad (4.80)$$

$$\text{où } \tilde{C}_{ij} = \frac{\chi_{ij} - g_{i \rightarrow j, j}(1 - \tanh^2 H_i)}{g_{j \rightarrow i, j}} + \tanh H_i \tanh H_j. \quad (4.81)$$

Un algorithme élémentaire d'inversion peut être décrit par les règles suivantes :

- $h_{i \rightarrow j} \leftarrow H_i - u_{j \rightarrow i}$
- Actualiser les messages $g_{i \rightarrow j, k}$ selon l'équation (4.75).
- Actualiser les messages J_{ij} selon (4.80).
- Actualiser les $u_{i \rightarrow j}$ selon (4.74).
- Actualiser les $v_{i \rightarrow j, k}$ selon (4.75).

Dans certaines situations, le graphe sous-jacent n'est pas connu *a priori*. Pour autant, la propagation inverse des susceptibilités *sur le graphe complet* est souvent capable d'inférer la structure graphique du modèle en faisant converger à zéro les couplages entre variables non voisines.

Nous avons notamment pu constater cette propriété sur la chaîne linéaire de spins, cf. (4.1). Partant d'une chaîne d'Ising quelconque, nous commençons par calculer exactement ses magnétisations locales et ses fonctions de corrélation à deux points (problème direct). Puis nous livrons ce jeu d'observables (χ_{ij}, H_i) à notre algorithme d'inférence inverse, qui opère sans connaissance *a priori* de l'ordre de la chaîne. L'algorithme est alors capable d'inférer la valeur exacte des couplages et des champs extérieurs (J_{ij}, h_i) , assignant la valeur 0 aux couplages J_{ij} de paires non voisines : il reconstruit ainsi l'ordre exact de la chaîne, ainsi que l'intensité des interactions reliant les variables consécutives, à partir des seules observables (χ_{ij}, H_i) . La reconstruction peut aussi être effectuée exactement quand le graphe est un arbre, pourvu que la dépendance aux conditions de bord (au niveau des branches) décroisse suffisamment vite avec le nombre de générations.

L'algorithme inverse de propagation des susceptibilités est en théorie applicable à n'importe quel modèle où l'approximation de Bethe fournit une bonne approximation. C'est le cas par exemple du modèle de Sherrington-Kirkpatrick (SK) [SK75], où les couplages J_{ij} sont tirés au hasard avec une loi normale de moyenne nulle et de variance J^2/N . Dans la phase de « haute température » $J < 1$, les méthodes de champ moyen approchent la solution réelle du problème *direct* $(J_{ij}, h_i) \rightarrow (\chi_{ij}, H_i)$ avec une précision arbitraire quand $N \rightarrow \infty$ [MP87]. Afin de tester l'algorithme in-

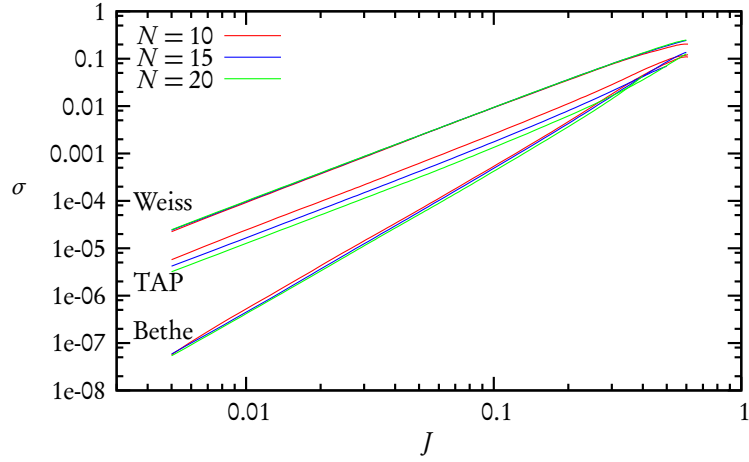


Fig. 4.7: Erreur moyenne $\sigma^2 = \frac{N}{J^2} \mathbb{E}[(J'_{ij} - J_{ij})^2]$ commise par les trois algorithmes présentés dans le texte. La propagation des convictions, qui repose sur l'approximation de Bethe, affiche les meilleures performances.

verse $(\chi_{ij}, H_i) \rightarrow (J_{ij}, h_i)$ sur ce modèle, une instance $\{J_{ij}\}$ de SK est tirée au hasard, et ses fonctions de corrélation sont estimées *de manière exhaustive*. Nous avons comparé trois algorithmes prenant en entrée ces fonctions de corrélation, et livrant en sortie une estimation $\{J'_{ij}\}$ des couplages : les deux premiers sont basés sur une approximation de champ moyen (Weiss et TAP), décrits dans [KR98], et le troisième est l'algorithme inverse de propagation des susceptibilités. La figure 4.7 compare les performances moyennes de ces trois algorithmes.

Malgré ses bonnes performances, la propagation des susceptibilités présente quelques inconvénients :

- Son exécution prend de l'ordre de N^3 opérations, contre N^2 pour le champ moyen. Elle requiert N^{n+1} opérations quand elle prend en compte les fonctions de corrélation d'ordre n .
- Son efficacité semble limitée aux phases « paramagnétiques », et est mise en échec par l'apparition d'ordre à longue portée. Ceci n'est guère surprenant, puisque ces limitations concernent également la résolution du problème direct.
- Les premiers essais d'application à des problèmes réels (non-aléatoires) tirés de [SBSB06], se sont avérés peu concluants.

Rien n'indique toutefois que cette dernière limitation soient inhérente à la méthode elle-même : il est probable que des versions raffinées, ou spécialement adaptées, puissent surmonter certaines des difficultés rencontrées.

Références

Le chapitre présent traite de la méthode de la cavité sous l'hypothèse de la symétrie des répliques, telle que développée par Mézard et Parisi [MP01] dans le contexte des graphes dilués. Comme nous l'avons déjà souligné, cette approche est équivalente à l'algorithme somme-produit décrit dans [KFL01]. Les articles de Yedidia *et al.* [Yed01, YFW02] établissent le lien entre propagation des convictions et physique statistique, et mettent en évidence la variationnalité de l'énergie libre.

Les performances de l'algorithme de propagation des convictions (BP) dans le contexte du décodage des codes LDPC a été étudiée avec beaucoup de détails dans [RU01, RSU01]. Auparavant, l'analyse de l'algorithme BP sur le canal d'effacement avait été effectuée à l'aide d'équations différentielles [LMS⁺97].

Le seuil optimal ϵ_c du canal d'effacement a originellement été calculé grâce à l'astuce des répliques [FL03]. Nous avons repris ici ce calcul dans le cadre de la méthode de la cavité, en accord avec l'approche adoptée dans [MR06a, MR06b]. Ce calcul fait l'objet d'une preuve rigoureuse [MMRU04], dans laquelle décodages itératif et optimal sont mis en rapport par le truchement d'une construction de Maxwell [MMU05] ; [RU07] reprend les éléments importants de cette preuve.

Les articles [MR06a, MR06b] entreprennent l'étude des grandes déviations des équations de cavité (4.46), (4.47) afin d'estimer la probabilité d'erreur du décodage optimal.

Le traitement du problème d'occupation généralise des travaux antérieurs [WH00, MMS06, ZM06, DMU04], où la méthode de la cavité a été appliquée à des cas particuliers de ce problème. L'exemple des sous-parties d'arrêt est quant à lui spécifique à cette thèse.

L'extension de la cavité aux fonctions de corrélations, et son application au problème de la machine de Boltzmann, est le fruit d'une collaboration avec Marc Mézard, et fera l'objet d'une publication future.

Chapitre 5

Spectres de distance

Ce chapitre aborde la question de l'organisation géométrique des solutions d'un problème de satisfactions de contraintes. Après quelques préliminaires sur les méthodes combinatoires, et l'introduction d'un outil nouveau, la x -satisfaisabilité, celles-ci sont mises à profit afin de prouver la fragmentation dans les problèmes k -XORSAT et k -SAT. La question de la relation précise de la x -satisfaisabilité avec les phénomènes de fragmentation et d'ergodicité est abordée. Enfin, les propriétés de distance des codes sont discutées en rapport avec les performances de décodage.

5.1 Préliminaires : un peu de combinatoire

Le chapitre précédent a été l'occasion d'introduire les méthodes de la physique statistique, dont l'objectif est l'évaluation fiable des moyennes gelées. Celles-ci sont les plus pertinentes à la fois du point de vue de la physique, où elles décrivent l'équilibre thermodynamique, et du point de vue de la complexité algorithmique moyenne, qui s'intéresse aux propriétés *typiques* des problèmes aléatoires. Néanmoins, il est souvent plus aisé de procéder au calcul de la moyenne recuite :

$$f_{\text{recuit}} = - \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E}(Z) \quad (5.1)$$

Cette moyenne fournit une borne sur la moyenne gelée de l'énergie libre, et s'en approche souvent de manière spectaculaire. En outre, elle s'offre la plupart du temps à un traitement analytique rigoureux.

5.1.1 Le calcul recuit

Afin d'illustrer les propriétés des moyennes recuites, nous nous appuyons dans un premier temps sur l'exemple des problèmes d'occupation décrits au paragraphe 3.3. Pour une partie $A \subset \mathbb{N}$ arbitraire, nous considérons le problème d'occupation défini par :

$$S \text{ est une } A\text{-partie ssi pour chaque facteur } a, \quad |\partial a \cap S| \in A, \quad (5.2)$$

sur un graphe factoriel aléatoire de distributions (L, R) . Le nombre de A -parties de taille $|S| = w$ est noté n_w ; le calcul recuit consiste à estimer sa moyenne $\mathbb{E}n_w$. Afin d'identifier les parties S , nous recourons aux fonctions polynômes caractéristiques : on appelle $\text{coef}(p(x), x^i)$ le coefficient d'ordre i dans le polynôme $p(x)$. Par exemple, le nombre de manières de choisir d voisins parmi tous les voisins d'un facteur a s'écrit : $\text{coef}((1+x)^{|\partial a|}, x^d) = \binom{|\partial a|}{d}$.

Le tirage du graphe factoriel s'effectue de la manière suivante : chaque variable et chaque facteur est muni d'un nombre de « jambes » égal au nombre de voisin qu'il est supposé avoir : $L(0)N$ variables n'ont aucune jambe, $L(1)N$ en ont une, etc. De la même façon, pour chaque k , $R(k)M$ facteurs sont dotés de k jambes. Le nombre total de jambes de part et d'autre vaut $\mathbb{E}(\ell)N = \mathbb{E}(k)M$. Un graphe aléatoire correspond à un choix d'appariement entre les jambes des variables d'un côté, et les jambes des facteurs de l'autre ; ces appariements engendrent les liens du graphe factoriel.

Soit S une sous-partie des variables $\{1, \dots, N\}$, et e le nombre total de jambes attachées à S . Le nombre total de manières de choisir e jambes parmi les jambes des facteurs vaut :

$$\text{coef} \left[\prod_{a=1}^M (1+x)^{|\partial a|}, x^e \right] = \binom{\mathbb{E}(k)M}{e} \quad (5.3)$$

Ici l'occurrence de « x » signale une jambe de a appariée à une jambe de S . Maintenant, le nombre de manières de choisir les e jambes *tout en respectant la contrainte d'occupation* s'exprime comme :

$$\text{coef} \left[\prod_{a=1}^M p_{A,|\partial a|}(x), x^e \right], \quad \text{avec} \quad p_{A,k}(x) = \sum_{\substack{0 \leq d \leq k \\ d \in A}} \binom{k}{d} x^d. \quad (5.4)$$

La probabilité que la partie S satisfasse la condition d'occupation vaut donc :

$$\binom{\mathbb{E}(k)M}{e}^{-1} \text{coef} \left[\prod_k p_{A,k}(x)^{R(k)M}, x^e \right]. \quad (5.5)$$

Par ailleurs, le nombre de manières de choisir une partie S de taille w ayant e

jambes vaut [DRU06] :

$$\text{coef} \left[\prod_{\ell} (1 + u y^{\ell})^{L(\ell)N}, u^w y^e \right], \quad (5.6)$$

où u indique la présence d'une variable de S , et y celle d'une jambe de S . Finalement :

$$\mathbb{E}(n_w) = \sum_{e=0}^{\mathbb{E}(k)M} \text{coef} \left[\prod_{\ell} (1 + u y^{\ell})^{L(\ell)N}, u^w y^e \right] \frac{\text{coef} \left[\prod_k p_{A,k}(x)^{R(k)M}, x^e \right]}{\binom{\mathbb{E}(k)M}{e}}. \quad (5.7)$$

La méthode du col permet d'estimer le comportement asymptotique des différents termes intervenant dans cette expression quand $N \rightarrow \infty$. Le second d'entre eux est dominé par l'exposant

$$\lim_{N \rightarrow \infty} \frac{1}{N} \log \text{coef} \left[\prod_k p_{A,k}(x)^{R(k)M}, x^{N\epsilon} \right] = \frac{\mathbb{E}(\ell)}{\mathbb{E}(k)} \sum_k R(k) \log p_{A,k}(\bar{x}) - \epsilon \log \bar{x}, \quad (5.8)$$

où $\bar{x}$ vérifie :

$$\epsilon = \frac{\mathbb{E}(\ell)}{\mathbb{E}(k)} \sum_k R(k) \frac{\bar{x} p'_{A,k}(\bar{x})}{p_{A,k}(\bar{x})}. \quad (5.9)$$

Le premier terme s'évalue à :

$$\lim_{N \rightarrow \infty} \frac{1}{N} \log \text{coef} \left[\prod_{\ell} (1 + u y^{\ell})^{L(\ell)N}, u^{\omega N} y^{\epsilon N} \right] = \sum_{\ell} L(\ell) \log(1 + \bar{u} \bar{y}^{\ell}) - \omega \log \bar{u} - \epsilon \log \bar{y} \quad (5.10)$$

où $\bar{u}$ et $\bar{y}$ satisfont les équations de col :

$$\omega = \sum_{\ell} L(\ell) \frac{\bar{u} \bar{y}^{\ell}}{1 + \bar{u} \bar{y}^{\ell}}, \quad \epsilon = \sum_{\ell} \ell L(\ell) \frac{\bar{u} \bar{y}^{\ell}}{1 + \bar{u} \bar{y}^{\ell}} \quad (5.11)$$

On obtient ainsi :

$$\begin{aligned} \Omega(\omega) = \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E}(n_{N\omega}) &= \sum_{\ell} L(\ell) \log(1 + \bar{u} \bar{y}^{\ell}) + \frac{\mathbb{E}(\ell)}{\mathbb{E}(k)} \sum_k R(k) \log p_{A,k}(\bar{x}) \\ &\quad - \mathbb{E}(\ell) \log(1 + \bar{x} \bar{y}) - \omega \log(\bar{u}) \end{aligned} \quad (5.12)$$

où ω est donné par (5.11). La variable ϵ est évacuée, et les équations auto-cohérentes sur $\bar{x}$ et $\bar{y}$ se réécrivent sous la forme :

$$\bar{x} = \sum_{\ell} \tilde{\lambda}(\ell) \bar{u} \bar{y}^{\ell}, \quad \tilde{\lambda}(\ell) = \frac{\lambda(\ell)/(1 + \bar{u} \bar{y}^{\ell+1})}{\sum_{\ell'} \lambda(\ell')/(1 + \bar{u} \bar{y}^{\ell'+1})}, \quad (5.13)$$

$$\bar{y} = \sum_k \tilde{\rho}(k) \frac{q_{A,k+1}(\bar{x})}{p_{A,k}(\bar{x})}, \quad \tilde{\rho}(k) = \frac{\rho(k) \frac{p_{A,k}(\bar{x})}{p_{A,k+1}(\bar{x})}}{\sum_{k'} \rho(k') \frac{p_{A,k'}(\bar{x})}{p_{A,k'+1}(\bar{x})}}, \quad (5.14)$$

où $q_{A,k} = p'_{A,k}/k$, et où on a utilisé $p_{A,k} - xq_{A,k} = p_{A,k-1}$. Le cas de la fonction d'énumération des sous-parties d'arrêt [OVZ05] est décrit par :

$$p_{A,k} = (1+x)^k - kx, \quad q_{A,k} = (1+x)^{k-1} - 1 \quad (5.15)$$

et celui de la fonction d'énumération des poids de mots de codes [DRU06] par :

$$p_{A,k} = \frac{1}{2} \left[(1+x)^k + (1-x)^k \right], \quad q_{A,k}(x) = \frac{1}{2} \left[(1+x)^{k-1} - (1-x)^{k-1} \right]. \quad (5.16)$$

5.1.2 Comparaison avec la moyenne gelée

Une comparaison rapide entre les équations (4.54) et (5.12), avec $\bar{u} = 2^{-\beta}$ permet de se convaincre de la similarité entre les équations recuites et gelées. Tout se passe comme si les équations recuites offraient une approximation *factorisée* des équations gelées, reposant sur l'invariance des messages $x_{i \rightarrow a}$ et $y_{a \rightarrow i}$.

La similarité entre moyennes recuite et gelée devient une identité dans le cas d'un graphe régulier (ℓ, k) , car l'Ansatz factorisé y devient exact. L'entropie Ω prend donc la même valeur dans les deux cas, cf. (4.58) :

$$\Omega(\omega) = (1-\ell)H(\omega) - \omega\ell \log(x) + \frac{\ell}{k} \log p_{A,k}(x), \quad \omega = x \frac{q_{A,k}(x)}{p_{A,k}(x)}. \quad (5.17)$$

Cette identité a déjà été observée séparément dans les principaux cas particuliers du problème d'occupation : sur le problème des dimères sur graphe [ZM06], éq. (7), sur le comptage des cycles [MS06a], éq. (31), et sur celui des mots de code [Con02, DMU04].

La précision de l'approximation recuite peut s'avérer néanmoins excellente, même pour un graphe irrégulier. En témoigne par exemple la fonction d'énumération des sous-parties d'arrêt du code irrégulier défini par (4.49), dont les moyennes recuites et gelées, représentées figure 4.5, sont presque indistinguables. Pour les sous-parties d'arrêt comme pour les mots de code, une telle précision n'est cependant observée que pour les ensembles interdisant les degrés de variable 0 ou 1.

5.1.3 Ensemble « lâche »

Dans les raisonnements décrits ci-dessus, nous avons implicitement adopté un modèle particulier de graphes aléatoires, où le nombre de nœuds (variables ou facteurs) de degré donné est fixé préalablement au tirage du graphe. Il est intéressant de remarquer que le modèle poissonien utilisé dans k -SAT aléatoire, par exemple, ne rentre pas tout-à-fait dans ce cadre. Dans cet ensemble, le degré de chaque variable

est tiré indépendamment selon une loi poissonnienne, conditionnellement à ce que la somme de ces degrés vaille Mk . Dans la limite $N \rightarrow \infty$, le nombre de variables de degré ℓ se concentre autour de sa valeur moyenne $Ne^{-k\alpha}(k\alpha)^\ell/\ell!$, et l'on retrouve l'ensemble décrit plus haut. Cependant, s'il est vrai que ces deux ensembles sont *typiquement* équivalents, ils diffèrent dans leur propriétés de grandes déviations, dont dépend la moyenne recuite. En effet, supposons que l'entropie $\Omega(\omega)$ suive une loi de grande déviation :

$$\mathbb{P}\left(\frac{1}{N} \log n_{N\omega} = \Omega\right) \asymp 2^{-NL_\omega(\Omega)}. \quad (5.18)$$

La moyenne recuite vaut :

$$\lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E}(n_{N\omega}) = \lim_{N \rightarrow \infty} \frac{1}{N} \log \int d\Omega 2^{N[\Omega - L_\omega(\Omega)]} = \max_{\Omega} [\Omega - L_\omega(\Omega)] \quad (5.19)$$

tandis que la moyenne gelée vaut $\operatorname{argmax}_{\Omega} L_\omega(\Omega)$. Ainsi, plus la fonction de grande déviation est « resserrée » autour de sa valeur typique, meilleure est l'approximation recuite. L'ensemble considéré au paragraphe précédent, en ce qu'il « force » les histogrammes de degrés à prendre leur valeur typique, est meilleur que celui où ils sont libres de fluctuer. Nous qualifions le premier type d'ensemble de *serré*, et le second de *lâche*.

Comment le calcul recuit se généralise-t-il à l'ensemble lâche ? Supposons pour simplifier que les degrés de variable suivent une loi $L(\ell)$, et que les facteurs ait degré constant¹ k . Tandis que le terme (5.5) reste inchangé, le facteur combinatoire (5.6) devient :

$$\mathbb{E}_{\ell_1, \ell_2, \dots, \ell_N} \left\{ \operatorname{coef} \left[\prod_{i=1}^N z^{\ell_i} (1 + u y^{\ell_i}), u^w y^e z^{Mk} \right] \right\} / \mathbb{P}_{\ell_1, \ell_2, \dots, \ell_N} \left(\sum_i \ell_i = Mk \right), \quad (5.20)$$

où les degrés de variables $\ell_i = |\partial i|$ sont tirés selon $L(\ell)$. La contribution asymptotique du dénominateur vaut 1 pourvu que $\mathbb{E}(\ell) = k\alpha$, où $\alpha = M/N$. Le logarithme du numérateur divisé par N tend vers :

$$\log [L(\bar{z}) + \bar{u}L(\bar{z}\bar{y})] - k\alpha \log \bar{z} - \omega \log \bar{u} - \epsilon \log \bar{y}, \quad (5.21)$$

avec pour équations de col :

$$k\alpha = \bar{z} \mathbb{E}(\ell) \frac{\lambda(\bar{z}) + \bar{u}\bar{y} \lambda(\bar{y}\bar{z})}{L(\bar{z}) + \bar{u}L(\bar{z}\bar{y})}, \quad \omega = \frac{\bar{u}L(\bar{z}\bar{y})}{L(\bar{z}) + \bar{u}L(\bar{z}\bar{y})}, \quad \epsilon = \frac{\bar{u}\bar{y} \bar{z} \mathbb{E}(\ell) \lambda(\bar{z}\bar{y})}{L(\bar{z}) + \bar{u}L(\bar{z}\bar{y})}. \quad (5.22)$$

¹La généralisation à une distribution générale de degrés de facteurs ne pose pas de difficulté supplémentaire. Il faut cependant préciser si le nombre total $|E|$ de liens est fixé à l'avance, ou bien s'il est laissé libre de fluctuer.

L'expression (5.12) de Ω est juste modifiée de telle sorte que le terme $\sum_{\ell} L(\ell) \log(1 + \bar{u} \bar{y}^{\ell})$ est remplacé par $\log[\bar{u} L(\bar{y} \bar{z}) + L(\bar{z})] - k \alpha \log \bar{z}$.

Ici, au contraire de l'ensemble serré, le cas poissonien $L(\ell) = e^{-k\alpha} (k\alpha)^{\ell} / \ell!$ se simplifie considérablement. On trouve

$$\bar{x} = \frac{\omega}{1-\omega}, \quad \bar{y} = \frac{q_{A,k}(\bar{x})}{p_{A,k-1}(\bar{x})}, \quad \bar{z} = \frac{1}{1-\omega(1-\bar{y})}, \quad \bar{u} = \frac{\omega}{1-\omega} e^{k\alpha \bar{z}(1-\bar{y})} \quad (5.23)$$

d'où l'on tire :

$$\Omega(\omega) = H(\omega) + \alpha \log \left[p_{A,k} \left(\frac{\omega}{1-\omega} \right) (1-\omega)^k \right]. \quad (5.24)$$

Ce résultat peut d'ailleurs être retrouvé directement et de manière beaucoup plus simple. En effet, dans l'ensemble poissonien, les facteurs sont tirés indépendamment, ce qui permet leur factorisation :

$$\mathbb{E}(n_w) = \mathbb{E} \left[\sum_{\sigma, \|\sigma\|=w} \prod_{a=1}^M \mathbb{I} \left(\sum_{i \in \partial a} \sigma_i \in A \right) \right] = \sum_{\sigma, \|\sigma\|=w} \mathbb{P}_J \left(\sum_{i \in J} \sigma_i \in A \right)^M \quad (5.25)$$

La variable σ représente une sous-partie $S \subset \{1, \dots, N\}$. J est une collection de k variables choisies au hasard formant le voisinage d'un facteur aléatoire. La probabilité que $\sum_{j \in J} \sigma_j \in A$ ne dépend en fait que de la taille $w = N\omega$ de S . Plus spécifiquement on a :

$$\mathbb{E}(n_w) = \binom{N}{N\omega} \left[\sum_{\substack{0 \leq d \leq k \\ d \in A}} \binom{k}{d} \omega^d (1-\omega)^{k-d} \right]^M. \quad (5.26)$$

L'estimation asymptotique de cette expression redonne (5.24).

En guise d'illustration, les moyennes recuites dans les ensembles serré et lâche sont comparées figure 5.1 sur l'exemple de la fonction d'énumération des poids ($A = 2\mathbb{N}$) de 5-XORSAT homogène. Notez que dans cet exemple, l'existence d'une phase fragmentée rend la propagation des convictions impropre à décrire la moyenne gelée. Le chapitre 6 expose les principes permettant d'effectuer ce calcul avec multiplicité d'états.

5.2 x -satisfaisabilité et fragmentation

Une instance de problème de satisfaction de contraintes est dit x -satisfaisable s'il admet une paire de solutions σ, σ' séparées par une distance de Hamming $\approx Nx$:

$$\sum_{i=1}^N \left(1 - \delta_{\sigma_i, \sigma'_i} \right) \in [Nx - \epsilon(N), Nx + \epsilon(N)] \quad (5.27)$$

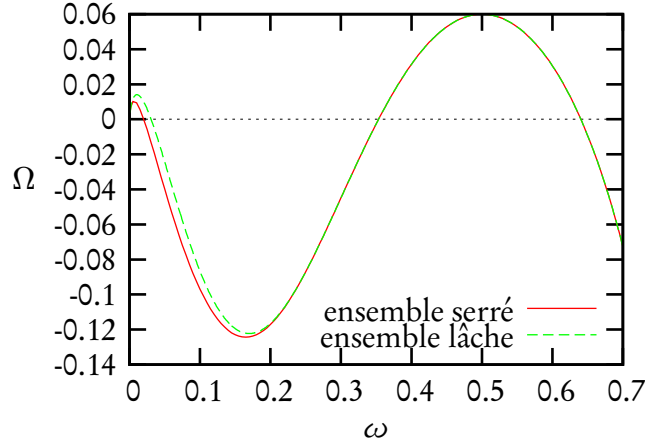


Fig. 5.1: Moyenne recuite de la fonction d'énumération des poids du problème 5-XORSAT avec $\alpha = M/N = 0,94$. Ω donne le taux de croissance du nombre moyen de solutions au problème homogène, (cf. (3.10) avec $\tau = 0$) en fonction du poids $\|\sigma\| = N\omega$ de σ . Les deux moyennes recuites majorent la moyenne gelée, qui décrit le comportement des instances typiques. Conformément à la discussion du texte, l'ensemble serré permet de mieux approcher la moyenne gelée.

où la résolution $\epsilon(N)$ est une fonction sous-extensive définie comme éq. (2.16), par exemple $\epsilon(N) = \sqrt{N}$. Dans la suite x sera appelée distance réduite, c'est-à-dire la distance renormalisée par N .

La plupart des problèmes subissant une transition SAT/non-SAT abrupte connaissent également une transition de x -satisfaisabilité. Il existe dans ce cas un seuil $\alpha_s(x)$ tel que, pour $N \rightarrow \infty$ avec $\alpha = M/N$ fixé :

- pour $\alpha = M/N < \alpha_s(x)$, une instance est x -satisfaisable presque sûrement,
- pour $\alpha = M/N > \alpha_s(x)$, elle est non- x -satisfaisable presque sûrement.

Dans l'article [MMZ05b] nous pensions détenir une preuve de l'existence d'un seuil abrupt non-uniforme, mais il se trouve que cette preuve est erronée. La version de l'article reproduit à la fin du texte doit remplacer celle présente sur le serveur de preprints : le seuil abrupt de la x -satisfaisabilité y est maintenant présenté comme une conjecture.

La notion de x -satisfaisabilité permet d'explorer le spectre des distances accessibles dans les problèmes de satisfaction de contrainte. Elle renseigne sur les propriétés géométriques de l'espace des solutions, et permet à ce titre d'étudier le phénomène de fragmentation décrit au paragraphe 2.3.2.

5.2.1 x -satisfaisabilité dans k -XORSAT

Dans le problème XORSAT, la question de l' x -satisfaisabilité est intimement liée à la fonction d'énumération des poids. Soit σ^0 une solution de référence à un problème booléen linéaire. Il est facile de voir que l'ensemble des solutions à distance w de σ^0 est isomorphe à l'ensemble des solutions de poids w au problème homogène ($\tau = 0$) : en effet, si σ' est une solution telle que $\sum_i |\sigma'_i - \sigma^0_i| = w$, alors $\sigma = \sigma' - \sigma^0$ est solution du problème homogène avec $\sum_i |\sigma_i| = w$, et réciproquement. Le spectre des distances du problème original est donc entièrement reflété par le spectre des poids du problème homogène.

Commençons par exploiter cette remarque en l'appliquant directement à l'ensemble k -XORSAT. En examinant la moyenne recuite de la fonction d'énumération des poids (Fig. 5.1), on observe un « fossé » de distances inaccessibles. Partout où l'entropie est négative, l'inégalité de Markov :

$$\mathbb{P}(n_{Nx} \geq 1) \leq \mathbb{E}(n_{Nx}) \asymp 2^{N\Omega(x)} \rightarrow 0 \quad (5.28)$$

implique l'impossibilité de trouver des solutions à distance $x (= \omega)$: l'instance est donc presque sûrement non- x -satisfaisable.

Ce constat est cohérent avec le scénario de fragmentation, qui prévoit que les solutions appartenant à des amas différents soient extensivement éloignées. Conformément à cette image, la figure 5.1 met en évidence l'existence de deux zones de distances possibles : on trouve des solutions près de σ^0 (appartenant au même amas), ou loin de σ^0 (provenant d'un amas différent), mais pas aux distances intermédiaires.

Le raisonnement précédent est cependant loin de constituer une preuve complète de la fragmentation. En particulier, il ne prouve pas que les zones mentionnées sont toutes les deux peuplées. Par ailleurs, le fossé disparaît pour certaines densités supérieures à α_d , et n'existe pas du tout pour $k = 3$.

Ces difficultés peuvent toutefois être surmontées grâce à l'analyse du processus d'effeuillage. Rappelons que pour $\alpha_d < \alpha < \alpha_c$, l'effeuillage laisse place à un cœur extensif. On peut montrer [CDMM03, MRTZ03] que ce cœur se comporte comme un graphe factoriel aléatoire ayant pour distribution de degrés une loi de Poisson tronquée :

$$L(\ell) = \frac{1}{e^\lambda - 1 - \lambda} \frac{\lambda^\ell}{\ell!} \mathbb{I}(\ell \geq 2), \quad (5.29)$$

où λ est solution de l'équation :

$$\lambda = k\alpha(1 - e^{-\lambda})^{k-1}. \quad (5.30)$$

Le nombre de variables impliquées dans le cœur vaut :

$$N_c = N \left[1 - (1 + \lambda)e^{-\lambda} \right]. \quad (5.31)$$

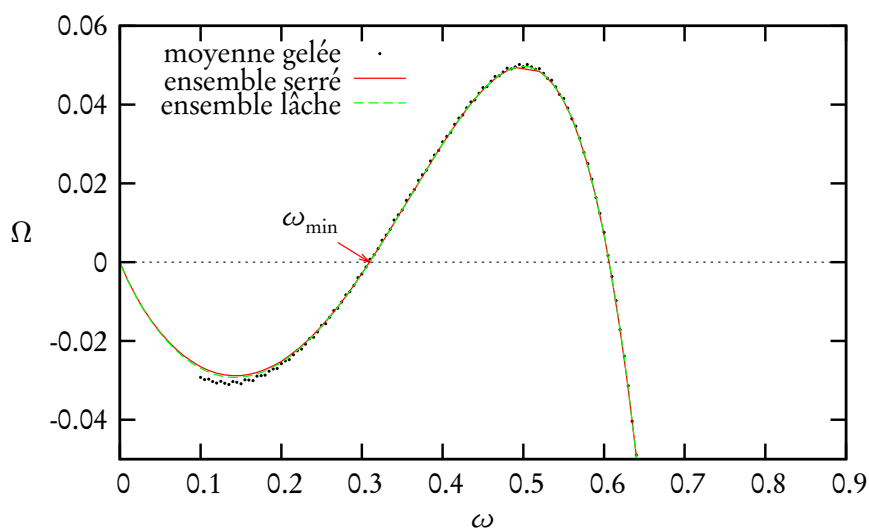


Fig. 5.2: Taux de croissance de la fonction d'énumération des poids dans le cœur de 3-XORSAT, pour $\alpha = 0,875$. Les moyennes gelées et recuites sont représentées. La distance réduite $\omega_{\min}$, où Ω s'annule, définit la distance minimale entre deux solutions du cœur. De même que la moyenne recuite majore la moyenne gelée, la distance où le taux recuit s'annule minore la distance minimale typique. Notez que la distance réduite ω est prise relativement à la taille du cœur N_c , et non pas à la taille totale N .

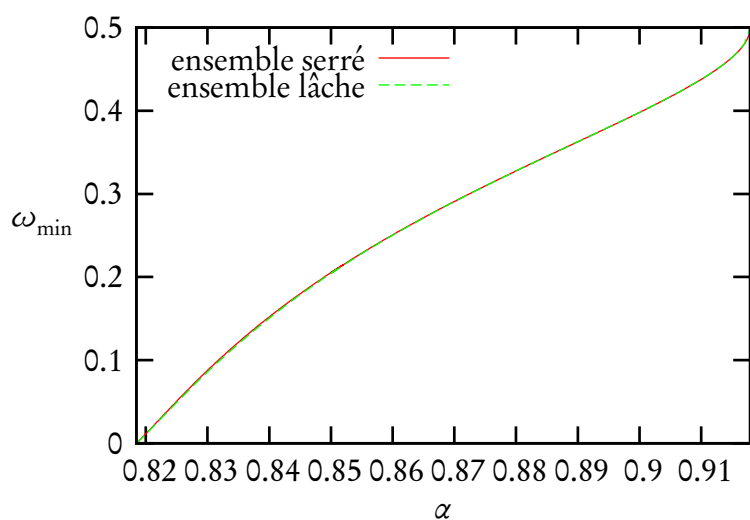


Fig. 5.3: Bornes inférieures sur la distance minimale $\omega_{\min}$ du cœur de 3-XORSAT, en fonction de la densité de tests $\alpha = M/N$.

La fonction d'énumération des poids de cet ensemble, dont les moyennes recuites et gelée sont représentées figure 5.2, est presque sûrement nulle sur un intervalle $]0, \omega_{\min}(\alpha)[$, où $\omega_{\min}(\alpha) > 0$ est la *distance minimale* du «code»² défini par le sous-problème de XORSAT restreint au cœur. Par conséquent, le cœur est non- x -satisfaisable presque sûrement pour tout $x \in]0, \omega_{\min}[$.

Les solutions du cœur sont donc séparées les unes des autres par une distance supérieure à $N_c \omega_{\min}$. La moyenne recuite fournit une borne inférieure à cette distance, cf. fig. 5.3. Ainsi, deux solutions du problème complet, construites à partir de deux solutions distinctes du cœur, sont *a fortiori* extensivement éloignées. La réciproque, selon laquelle deux solutions complètes issues de la même solution de cœur soit joignables par des sauts d'au plus $\epsilon(N)$ variables, est démontrée dans [MRTZ03]. La combinaison de ces deux preuves valide *a posteriori* la définition des amas comme ensemble de solutions issues d'une même solution de cœur (§3.2.1).

D'une certaine manière, l'espace des solutions du cœur de XORSAT est lui-même fragmenté, puisque les solutions sont extensivement séparées les unes des autres. Cependant, au contraire du problème XORSAT complet, les «amas» du cœur ne connaissent pas ici de fluctuations internes, car ils sont constitués d'une unique solution. La méthode de la cavité à un état unique, que nous avons utilisée pour calculer la moyenne gelée dans la figure 5.2, permet d'étudier la statistique de ces amas-singletons, et reste valide. La présence de fragmentation ne suffit donc pas à violer l'hypothèse d'un état unique : cette violation requiert en plus l'imbrication de plusieurs niveaux de fluctuations.

Cette remarque vaut également pour les mots de codes et les sous-parties d'arrêt des codes linéaires dilués. Bien qu'extensivement éloignés, les solutions de ces problèmes sont appréhendables par un état unique, ainsi que nous l'avons implicitement supposé dans les paragraphes 4.2.1 et 4.2.2.

5.2.2 x -satisfaisabilité dans k -SAT

L'existence d'une définition constructive des amas de XORSAT facilite considérablement l'analyse de la fragmentation dans ce problème. Dans le cas de la k -satisfaisabilité en revanche, l'absence de caractérisation des amas nous contraint à nous reposer uniquement sur la notion de x -satisfaisabilité. De la même manière que nous avons utilisé la fonction d'énumération des poids pour étudier XORSAT, nous intro-

²Plus généralement, pour une construction aléatoire de codes LDPC, on peut montrer qu'une distance minimale existe, et donc que les mots de code sont bien séparés, dès que $\lambda'(0)\rho'(1) < 1$ [DRU06].

duisons ici pour chaque instance une fonction d'énumération des *distances* :

$$Z(x) = \sum_{\sigma, \tau} \mathbb{I} \left(\sum_{i=1}^N |\sigma_i - \tau_i| \in I_N(x) \right) \prod_{a=1}^M [\mathbb{I}(\sigma \models a) \mathbb{I}(\tau \models a)], \quad (5.32)$$

où $I_N(x) = [Nx - \epsilon(N), Nx + \epsilon(N)]$. Cette fonction est supérieure à 1 si et seulement si la formule est x -satisfaisable.

Par la suite nous calculerons deux bornes sur la x -satisfaisabilité : une borne supérieure obtenue à l'aide de la méthode dite du premier moment, et une borne inférieure basée sur la méthode du second moment. Ces deux bornes rigoureuses nous permettront de tirer des conséquences fermes sur la structure géométrique de l'espace des solutions.

La première partie de notre raisonnement consiste à utiliser la moyenne recuite de $Z(x)$ et l'inégalité de Markov pour identifier une zone du diagramme (α, x) où les formules sont presque sûrement non- x -satisfaisable. Cette méthode du « premier moment » reprend exactement celle du paragraphe 2.3.1. On a :

$$\begin{aligned} \mathbb{E}[Z(x)] &= \sum_{\sigma, \tau} \mathbb{I}[\|\sigma - \tau\| \in I_N(x)] \mathbb{E} \left\{ \prod_{a=1}^M [\mathbb{I}(\sigma \models a) \mathbb{I}(\tau \models a)] \right\} \\ &= 2^N \sum_{d \in I_N(x)} \binom{N}{d} \mathbb{P}(\sigma, \tau \models a \mid d = \|\sigma - \tau\|)^M \end{aligned} \quad (5.33)$$

où on a utilisé le fait que la probabilité sur les clauses se factorise. La probabilité qu'une clause aléatoire soit satisfaite par deux configurations ne dépend que de leur distance d :

$$\mathbb{P}(\sigma, \tau \models a \mid d = \|\sigma - \tau\|) = 1 - 2^{1-k} + 2^{-k}(1 - d/N)^k \quad (5.34)$$

En effet, pour chaque clause, seules deux des choix de négations sont interdits parmi 2^k . Néanmoins, si σ et τ coïncident sur les variables de a , ce qui se produit avec probabilité $(1 - d/N)^k$, un seul choix est interdit. Finalement, on a :

$$\mathbb{E}[Z(x)] \asymp 2^N [1 + H(x) + \alpha \log(1 - 2^{1-k} + 2^{-k}(1-x)^k)] \quad (5.35)$$

Ainsi, dès que :

$$\alpha > \alpha_{BS}(x) \doteq - \frac{1 + H(x)}{\log(1 - 2^{1-k} + 2^{-k}(1-x)^k)} \quad (5.36)$$

la formule est non- x -satisfaisable presque sûrement. $\alpha_{BS}(x)$ est une *borne supérieure* sur le seuil de x -satisfaisabilité $\alpha_s(x)$. Comme pour α_s , cette borne peut être améliorée par des techniques telles que celles décrites dans [DB97, KKKS98], mais le bénéfice reste quantitativement faible en pratique.

La seconde partie de notre raisonnement se fonde sur la méthode du « second moment », originellement développée par [AM02, AP04, ANP05] afin de minorer α_s . Cette méthode repose sur l'inégalité suivante. Pour toute variable Z positive ou nulle,

$$\mathbb{P}(Z > 0) \geq \frac{\mathbb{E}(Z)^2}{\mathbb{E}(Z^2)} \quad (5.37)$$

Naturellement, on a toujours $\mathbb{E}(Z^2) \geq \mathbb{E}(Z)^2$. Si l'on choisit $Z = Z(x)$ défini par l'équation (5.33), le rapport $\mathbb{E}(Z)^2/\mathbb{E}(Z^2)$ devient même exponentiellement petit en N , rendant l'inégalité inutilisable. La raison intuitive de cet échec tient à la forte proportion de littéraux satisfaits dans les paires de solutions dominant $Z(x)^2$. Or de telles paires sont très corrélées entre elles, car elles ont tendance à suivre la « règle de la majorité » en satisfaisant le plus de littéraux possibles. Ces corrélations sont la cause de l'échec de la méthode du second moment, qui repose sur la commensurabilité de $\mathbb{E}(Z^2)$ et de $\mathbb{E}(Z)^2$. Afin de rééquilibrer la mesure, de telle sorte que dominent les solutions ayant la moitié de leurs littéraux satisfaits, nous redéfinissons $Z(x)$ en introduisant des poids :

$$Z(x) = \sum_{\sigma, \tau} \mathbb{I}(\|\sigma - \tau\| = \lfloor Nx \rfloor) \prod_{a=1}^M W(\sigma, \tau, a) \quad (5.38)$$

où $W(\sigma, \tau, a)$ est une fonction positive ou nulle vérifiant : $W = 0$ ssi $\sigma \not\models a$ ou $\tau \not\models a$. Nous supposons en plus que $W(\sigma, \tau, a)$ ne dépend que :

- du nombre $s_a(\sigma)$ de littéraux de a satisfaits par σ ,
- du nombre $s_a(\tau)$ de littéraux de a satisfaits par τ ,
- du nombre $q_a(\sigma, \tau)$ de littéraux de a prenant la même valeur dans σ et τ .

Un choix simple, quoique pas nécessairement optimal, est :

$$W(\sigma, \tau, a) = \begin{cases} 0 & \text{si } s_a(\sigma) = 0 \text{ ou } s_a(\tau) = 0 \\ \lambda^{s_a(\sigma) + s_a(\tau) - q_a(\sigma, \tau)} & \text{sinon} \end{cases} \quad (5.39)$$

Les paramètres λ et ν devront être choisis de telle sorte que les paires échantillonnées par la moyenne de $Z(x)^2$ soient décorréliées.

Le premier moment de $Z(x)$ se calcule de la même manière que précédemment :

$$\mathbb{E}[Z(x)] = 2^N \binom{N}{\lfloor Nx \rfloor} f_1(x)^M \quad (5.40)$$

$$\begin{aligned} f_1(x) &= \mathbb{E}[W(\sigma, \tau, a)] = \sum_{\substack{n_1+n_2+n_3+n_4=k \\ n_1+n_2>0, n_1+n_3>0}} \frac{k!}{n_1!n_2!n_3!n_4!} \lambda^{2n_1+n_2+n_3} \nu^{n_1+n_4} x^{n_2+n_3} (1-x)^{n_1+n_4} \\ &= 2^{-k} ((1-x)\nu(1+\lambda^2) + 2x\lambda)^k - 2^{1-k} (x\lambda + (1-x)\nu)^k + 2^{-k} ((1-x)\nu)^k. \end{aligned} \quad (5.41)$$

Le calcul du second moment se fait selon les mêmes principes :

$$\mathbb{E} [Z(x)^2] = \sum_{\sigma, \tau, \sigma', \tau'} f_2(\mathbf{a})^M \quad (5.42)$$

où $f_2(\mathbf{a}) = \mathbb{E} [W(\sigma, \tau, a)W(\sigma', \tau', a)]$ ne dépend que du vecteur $\mathbf{a}$ des recouvrements entre les chaînes σ, τ, σ' et τ' . Pour chacune des huit possibilités de recouvrement entre quatre variables binaires, indicées par $u \in \{0, 1\}^3$, a_u est la proportion de variables i telles $(\tau_i - \sigma_i, \sigma'_i - \sigma_i, \tau'_i - \sigma_i) = u$. Par souci de concision, nous ne reproduisons pas ici l'expression de f_2 , et référons le lecteur à l'article [MMZ05b].

Le vecteur des recouvrements $\mathbf{a}$ étant fixé, le nombre de quadruplets $(\sigma, \tau, \sigma', \tau')$ réalisant ce motif vaut : $2^N N! / \prod_u (Na_u)!$. On a donc :

$$\begin{aligned} \mathbb{E} [Z(x)^2] &= 2^N \sum_{\mathbf{a} \in V} \frac{N!}{\prod_u (Na_u)!} f_2(\mathbf{a})^M \\ &\sim C_0 N^{3/2} \exp \left\{ N \ln(2) \max_{\mathbf{a} \in V} [1 + H_8(\mathbf{a}) + \alpha \log f_2(\mathbf{a})] \right\} \end{aligned} \quad (5.43)$$

où $H_8(\mathbf{a}) = -\sum_u a_u \log a_u$, et où V désigne le simplexe :

$$a_{100} + a_{101} + a_{110} + a_{111} = x, \quad a_{001} + a_{010} + a_{101} + a_{110} = x, \quad \sum_{v \in \{0,1\}^3} a_v = 1 \quad (5.44)$$

Voyons maintenant comment « équilibrer » cette somme. On veut que le maximum sur $\mathbf{a}$ soit atteint pour des paires décorréées, c'est-à-dire pour un recouvrement :

$$a_{000}^* = a_{001}^* = \frac{(1-x)^2}{2}, \quad a_{001}^* = a_{010}^* = a_{100}^* = a_{111}^* = \frac{x(1-x)}{2}, \quad a_{101}^* = a_{110}^* = \frac{x^2}{2} \quad (5.45)$$

Pour cette valeur, l'entropie H_8 est maximale sur V et vaut $1 + 2H(x)$. Par ailleurs, on a $f_2(\mathbf{a}^*) = f_1(x)^2$.

Afin que $\mathbf{a}^*$ soit le maximum de $H_8 + \alpha \log f_2$, on impose la condition nécessaire³ $\partial_{\mathbf{a}} f_2(\mathbf{a}^*) = 0$ qui s'écrit, avec le choix (5.39) :

$$\begin{aligned} [\nu(1-x)]^{k-1} &= (\lambda^2 + 1 - 2\lambda\nu)(2\lambda x + \nu(1-x)(1+\lambda^2))^{k-1} \\ (\nu(1-x) + \lambda x)^{k-1} &= (1 - \lambda\nu)(2\lambda x + \nu(1-x)(1+\lambda^2))^{k-1}. \end{aligned} \quad (5.46)$$

Ce choix des paramètres λ et ν étant fait, les estimations asymptotiques des premier et second moments de $Z(x)$ permettent de minorer le rapport :

$$\frac{\mathbb{E}(Z)^2}{\mathbb{E}(Z^2)} \geq C_1 \exp \left\{ N \ln(2) \min_{\mathbf{a} \in V} [1 + 2H(x) - H_8(\mathbf{a}) + 2\alpha \log f_1(x) - \alpha \log f_2(\mathbf{a})] \right\}. \quad (5.47)$$

³La dérivée $\partial_{\mathbf{a}}$ est prise le long du simplexe V

Pour peu que le minimum soit effectivement atteint en $\mathbf{a}^*$, le terme de l'exponentielle s'annule et $\mathbb{E}(Z)^2/\mathbb{E}(Z^2)$ est minoré par la constante $C_1 > 0$ quand $N \rightarrow \infty$. Sous l'hypothèse d'un seuil abrupt, cette minoration suffirait à assurer la x -satisfaisabilité presque sûrement. En l'état, la x -satisfaisabilité est démontrée *avec probabilité finie*.

La condition $\partial_{\mathbf{a}} f_2(\mathbf{a}^*)$ assure bien la stationnarité de la fonction au point de dé-correlation, mais pas sa minimalité. Cette dernière est néanmoins réalisée dès que :

$$\alpha < \alpha_{BI}(x) \doteq \inf_{\mathbf{a} \in V} \frac{1 + 2H(x) - H_8(\mathbf{a})}{\log f_2(\mathbf{a}) - 2 \log f_1(x)}. \quad (5.48)$$

La borne inférieure $\alpha_{BI}(x)$ ainsi obtenue est représentée figure 5.4 pour $k = 8$, accompagnée de la borne supérieure $\alpha_{BS}(x)$. Si l'on prend une tranche horizontale de ce diagramme (α fixé), on retrouve, pour $k \geq 8$, l'image déjà suggérée par la figure 5.1 dans le cas de XORSAT : on a ainsi démontré que les solutions s'organisent en deux régions disjointes du spectre des distances, laissant une zone de distances interdites entre elles : la première région, correspondant aux petites distances, contient les paires de solutions de même amas, tandis que la seconde, située autour de $x = 1/2$, contient les paires de solutions d'amas différents. Un fossé sépare ces deux régions, qui correspond à la zone interdite séparant les amas. La propriété ainsi démontrée suffit à valider l'hypothèse de la fragmentation, bien qu'elle n'en soit pas nécessairement consécutive.

Quand k devient grand, les estimations numériques, appuyées par des heuristiques analytiques sur la localisation de l'infremum de (5.48), conduisent à proposer la conjecture suivante :

$$\text{pour tout } x > 0, \quad \alpha_s(x) \sim 2^k \ln(2) \frac{1 + H(x)}{2} \quad (5.49)$$

avec convergence uniforme sur tout intervalle $[x_0, 1]$, $x_0 > 0$. D'autre part, on a [AP04]

$$\alpha_s(0) = \alpha_s \sim 2^k \ln(2). \quad (5.50)$$

Ainsi, le seuil critique $\alpha_{\text{fos}}(k)$, à partir duquel un fossé de distances inaccessibles apparaît, se comporte comme $2^{k-1} \ln(2)$ quand k tend vers l'infini. Ce seuil doit être comparé au seuil dynamique $\alpha_d(k) \sim 2^k \ln(k)/k$ prévu par la physique statistique, au dessus duquel la mesure se fragmente.

L'échec de notre méthode pour $k < 8$ peut avoir deux sources. D'abord, nous travaillons avec des bornes, ce qui réduit la précision de l'analyse. Ensuite, le seuil de fossé $\alpha_{\text{fos}}(k)$ ne correspond pas forcément au seuil de fragmentation α_d , comme nous allons le voir.

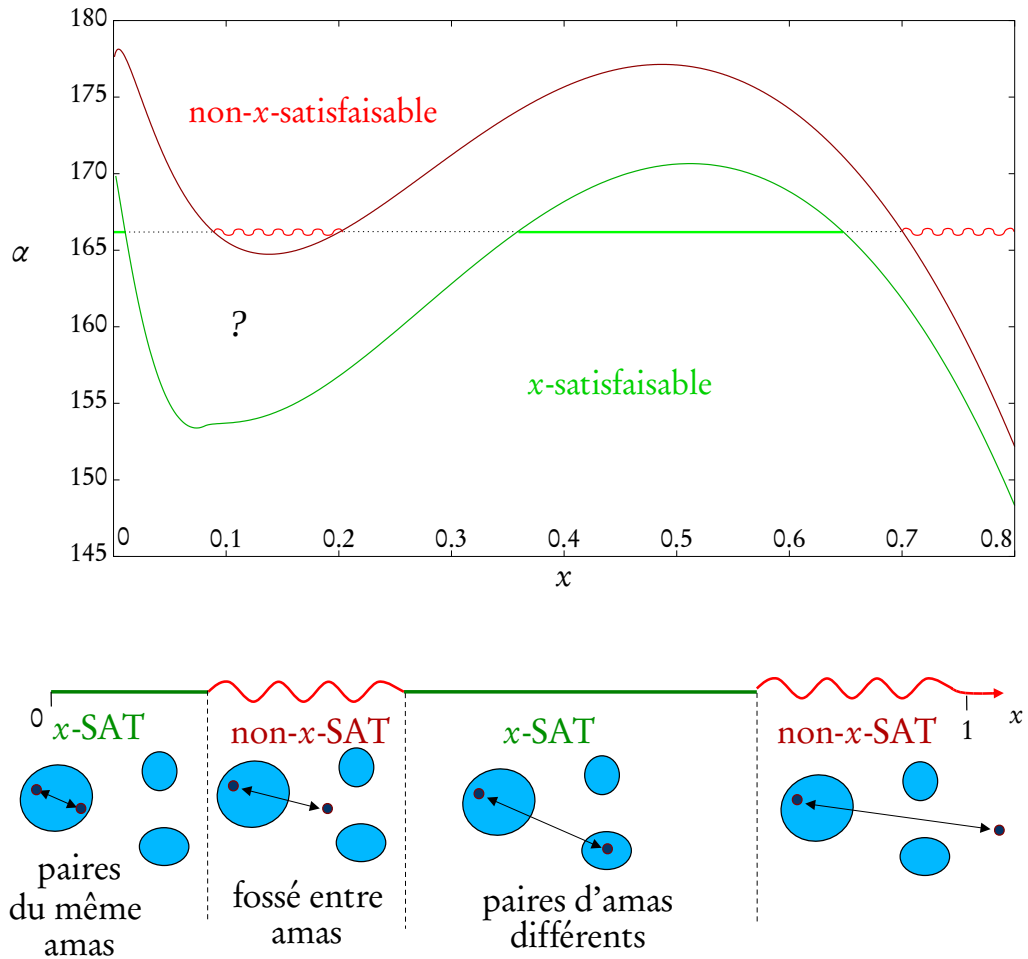


Fig. 5.4: Figure du haut : Bornes inférieure et supérieure sur le seuil de x -satisfaisabilité pour le problème 8-SAT aléatoire. À α fixé, le spectre des distances accessibles se divise en deux régions (schéma du bas) : soit les paires de solutions sont proches, et elles appartiennent au même amas ; soit elles sont éloignées, et dans ce cas elle appartiennent à des amas distincts. Entre ces deux régions se trouve un « fossé » de distances inaccessibles.

5.2.3 L' x -satisfaisabilité dans le modèle à amas aléatoires

Le modèle jouet introduit au paragraphe 2.3.3 permet de mettre en lumière, dans un cadre bien contrôlé, les limites de l' x -satisfaisabilité, et notamment les différences observées entre les seuils $\alpha_{\text{fos}}(k)$ et $\alpha_d(k)$.

L'estimation du seuil de x -satisfaisabilité dans le modèle à amas aléatoires passe par le calcul de trois quantités :

- le diamètre maximal d'un amas, c'est-à-dire la distance maximale entre deux solutions d'une même amas.
- la distance minimale entre amas, c'est-à-dire la distance minimale entre deux solutions d'amas distincts.
- la distance maximale entre deux solutions d'amas distincts.

La première quantité s'obtient en remarquant que le diamètre d'un amas vaut exactement le nombre de variables non gelées dans cet amas. Par conséquent, le diamètre (réduit) maximal d'un amas vaut $x_1 = s_M$, où s_M est la plus grande racine de $1 - \alpha - D(s||1 - p)$.

Considérons deux amas A et B pris au hasard, et examinons la probabilité que leur distance (réduite) minimale vaille x . Cette distance est donnée par le nombre de variables gelées i telles que $\pi_A(i) \neq \pi_B(i)$, ce qui est vérifié pour chaque variable avec probabilité $p^2/2$. La distance entre A et B suit donc une loi binomiale de paramètre $p^2/2$, et le nombre moyen $n(x)$ de couples d'amas à distance Nx se concentre presque sûrement autour de :

$$\mathbb{E}[n(x)] = 2^{2(1-\alpha)} \binom{N}{Nx} \left(1 - \frac{p^2}{2}\right)^{(1-x)N} \left(\frac{p^2}{2}\right)^{xN} \asymp 2^{Ns_2(x)}, \quad (5.51)$$

quand $s_2(x) \doteq 2(1-\alpha) - D(x||p^2/2)$ est strictement positif, et est presque sûrement nul sinon. Ainsi, la distance réduite minimale x_2 entre deux solutions d'amas distincts est définie comme la plus petite racine de $s_2(x)$. Un raisonnement très semblable donne la distance maximale entre deux amas : $x_3 = 1 - x_2$.

Le seuil d' x -satisfaisabilité vaut ainsi (cf. figure 5.5) :

$$\alpha_s(x) = \begin{cases} 1 & \text{si } x \in [0, 1 - p] \cup [p^2/2, 1 - p^2/2] \\ 1 - D(x||1 - p) & \text{si } x \in [1 - p, x_0] \\ 1 - \frac{1}{2}D(x||p^2/2) & \text{si } x \in [x_0, p^2/2] \\ 1 - \frac{1}{2}D(1 - x||p^2/2) & \text{si } x \in [1 - p^2/2, 1] \end{cases} \quad (5.52)$$

où x_0 est racine de $D(x||p^2/2) = D(x||1 - p)$.

Ce calcul appelle plusieurs remarques. Premièrement, il apparaît maintenant tout-à-fait normal que le seuil α_{fos} ne coïncide pas avec le seuil de *séparabilité*, noté

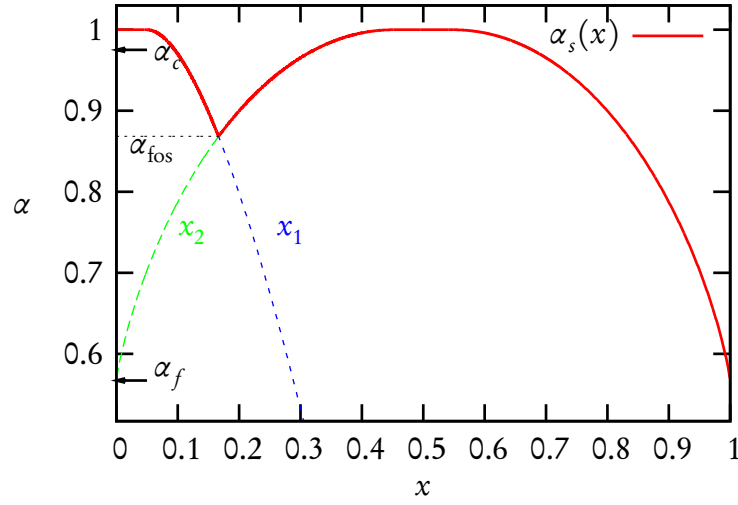


Fig. 5.5: Seuil de x -satisfaisabilité dans le modèle à amas aléatoires, pour $p = 0,95$. Les seuils de condensation α_c , de fossé α_{fos} , et de séparabilité α_f , sont représentés.

α_f , au dessus duquel les amas sont extensivement séparés deux-à-deux. En effet, pour $\alpha \in [\alpha_f, \alpha_{\text{fos}}]$, le diamètre maximal d'un amas est plus grand que la distance minimale entre deux amas : le spectre de x -satisfaisabilité voit donc un continuum de distances possibles, sans pouvoir détecter la fragmentation.

Mais le seuil α_f , au delà duquel les amas sont bien séparés, correspond-il lui-même à l'idée qu'on se fait du seuil « dynamique » α_d signalant la brisure d'ergodicité ? En effet, même pour $\alpha < \alpha_f$, il n'est pas à exclure que la majorité des amas continuent à être bien séparés les uns des autres, ou encore, quand bien même ils ne le seraient pas, que le saut d'un amas à l'autre reste très improbable.

Pour une configuration σ donnée, considérons le nombre d'amas aléatoires auxquels cette configuration appartient. Ce nombre est donné par une loi binomiale, et vaut en moyenne :

$$n = 2^{(1-\alpha)N} \left(1 - \frac{p}{2}\right)^N \quad (5.53)$$

Pour $\alpha > \alpha_d \doteq \log(2 - p)$, $n = 0$ presque sûrement. Autrement dit, les solutions sont exponentiellement rares. Pour $\alpha < \alpha_d$ au contraire, presque toute configuration est solution, et l'espace est trivialement ergodique. Ce seuil α_d correspond également à la densité pour laquelle $s_{\text{tot}} = 1 - \alpha + \log(2 - p)$ atteint sa valeur théorique maximale $s_{\text{tot}} = 1$.

La question demeure de savoir si l'espace des solutions peut être non-ergodique pour $\alpha \in [\alpha_d, \alpha_f]$, malgré l'absence de séparabilité. Pour y répondre, nous analysons une marche aléatoire uniforme sur l'espace des solutions. Soit A un amas d'entropie

interne s . Prenant une solution de cet amas, quelle est la probabilité qu'après t pas, le marcheur se trouve dans un amas d'entropie interne s' , noté B ? Soit a la proportion de variables à la fois libres dans A , et gelées dans B . La probabilité de a est donnée par le nombre total de manières de partitionner les N variables entre quatre groupes (gelées dans A et B , gelées dans A mais pas dans B , gelées dans B mais pas dans A , gelées ni dans A ni dans B), divisé par le nombre total de manières de partitionner les variables (gelées ou libres), indépendamment dans A et B :

$$q(a) = \frac{1}{\binom{N}{N_s} \binom{N}{N_{s'}}} \frac{N!}{(Na)! [N(s-a)]! [N(1-s'-a)]! [N(s'-s-a)]!}. \quad (5.54)$$

Pour que le marcheur tombe dans B , il faut que son évolution l'amène à prendre exactement la valeur requise par B sur ces Na variables. Après un temps $t \sim \beta N$, la probabilité que cela se produise vaut $\approx t 2^{-aN}$. Par ailleurs, les variables qui sont gelées à la fois dans A et B doivent coïncider, ce qui se produit avec probabilité $2^{N(s'+a-1)}$. Finalement, le taux de transition entre un amas de taille s et un autre de taille s' est majoré par (borne d'union) :

$$\phi(s \rightarrow s') \leq 2^{N\Sigma(s')} \sum_a q(a) t 2^{-aN} 2^{N(s'+a-1)} \quad (5.55)$$

Le maximum de $q(a)$ valant asymptotiquement 1 à l'ordre exponentiel, on a

$$\frac{1}{N} \log \phi(s \rightarrow s') \leq \Sigma(s') + s' - 1 \quad (5.56)$$

Tant que $\alpha > \alpha_d$, cette quantité est toujours négative, et les sauts d'amas à amas restent exceptionnels. Il y a bien brisure d'ergodicité.

Il faudrait nuancer cette analyse en notant que l'absence de notion d'énergie impose de se cantonner à l'espace strict des solutions, privant ainsi le marcheur des ponts de basse énergie pourtant caractéristiques des modèles réels. Le modèle à amas aléatoires a ceci de particulier qu'il construit explicitement l'espace des solutions comme une partition d'amas. En conséquence, il ne connaît pas de phase liquide non-triviale : $\alpha = \alpha_d$ correspond à une densité de contraintes nulle, car toute configuration est y solution. En dépit de cette simplification certes abusive, le modèle permet d'éclairer les différences entre les notions de séparabilité, d'ergodicité et d' x -satisfaisabilité.

5.3 Distances et erreur dans les codes linéaires

Les propriétés de distances des codes linéaires interviennent de manière cruciale dans les performances de décodage. La structure géométrique de l'environnement

d'un mot de code transmis nous a déjà servi de base au chapitre 1 (§1.2.2) dans l'estimation de la performance optimale des codes aléatoires.

Nous analysons ici les propriétés d'erreur de codes linéaire dilués (LDPC) sur le canal d'effacement (BEC). Deux types de décodage sont étudiés, pour lesquels deux types de fonctions d'énumération des A -parties sont exploitées afin de dériver des bornes rigoureuses sur la probabilité d'erreur moyenne :

- Le décodage optimal, qui repose sur le spectre de distances des mots de codes ($A = 2\mathbb{N}$).
- Le décodage itératif par propagation des convictions (§4.2.1), dont le succès dépend du spectre des tailles des sous-parties d'arrêt ($A = \mathbb{N} \setminus \{1\}$).

5.3.1 Ensemble expurgé

Au chapitre 1 nous avons indiqué que la probabilité d'erreur des meilleurs codes décroissent exponentiellement quand $N \rightarrow \infty$. Cependant, ainsi que l'a noté Gallager, cette décroissance s'avère être polynomiale en moyenne dans le cas des codes LDPC. Ce comportement est gouverné par une minorité de codes, en proportion polynomialement faible, pour lesquels il existe un mot de code très proche de $(0, \dots, 0)$, provoquant ainsi des erreurs avec probabilité finie. La solution préconisée par Gallager consiste à « débarasser » l'ensemble de ces mauvais codes, avec pour résultat un ensemble expurgé de codes ayant de bonnes propriétés de distance.

D'un point de vue technique, l'ensemble expurgé s'obtient en imposant une coupure dans le spectre des tailles $|S|$ de A -parties (pour les mots de code comme pour les sous-parties d'arrêt) : dès qu'un graphe admet une A -partie de taille $N\omega$ pour laquelle le taux de croissance $\Omega(\omega)$ est négatif, ce graphe est retiré de l'ensemble. On se débarrasse ainsi des distances atypiques.

5.3.2 Bornes d'union

Les événements conduisant à l'erreur du décodage optimal ou itératif peuvent s'expliquer à l'aide d'un formalisme commun. On peut supposer, sans perte de généralité, que le mot de code trivial $(0, \dots, 0)$ est transmis. Lors du passage dans le canal d'effacement, certains bits sont correctement transmis. Du point de vue graphique, cela revient à supprimer du graphe factoriel les variables correspondantes. Le décodage optimal échouera si le graphe restant admet une A -partie non-triviale, avec $A = 2\mathbb{N}$. Idem pour le décodage itératif, avec $A = \mathbb{N} \setminus \{1\}$.

Appelons E l'ensemble des bits effacés. Nous majorons la probabilité que cet

ensemble contienne une A -partie par la *borne d'union* suivante :

$$\begin{aligned} P_{\text{err}}(E) &= \mathbb{P}(\exists S \subset E \mid S \text{ est une } A\text{-partie}) \\ &\leq \min \left[\sum_{S \subset E} \mathbb{P}(S \text{ est une } A\text{-partie}), 1 \right] \end{aligned} \quad (5.57)$$

Cette probabilité est prise à la fois par rapport au choix aléatoire des bits effacés, et à celui du code expurgé. Fixons la taille de S à w . La probabilité qu'une partie S aléatoire de taille w soit une A -partie vaut $\mathbb{E}_{\text{exp}}(n_w) \binom{N}{w}^{-1}$, où la moyenne est prise dans l'ensemble expurgé. La probabilité d'erreur est alors majorée par

$$\begin{aligned} \mathbb{E}_{\mathcal{C}} [P_e(\mathcal{C})] &= \sum_E \epsilon^{|E|} (1 - \epsilon)^{N - |E|} P_{\text{err}}(E) \\ &\leq \sum_{|E|=0}^N \binom{N}{|E|} \epsilon^{|E|} (1 - \epsilon)^{N - |E|} \min \left[\sum_{w=0}^{|E|} \binom{|E|}{w} \mathbb{E}_{\text{exp}}(n_w) \binom{N}{w}^{-1}, 1 \right] \end{aligned} \quad (5.58)$$

et l'exposant minoré par :

$$E_{\text{expurg}} \geq E_{BU} \doteq - \max_{e \in [0,1]} \left\{ -D(e \parallel \epsilon) + \min \left[\max_{\omega \in [\omega_{\min}, e]} (\Omega(\omega) + eH(\omega/e) - H(\omega)), 0 \right] \right\} \quad (5.59)$$

Cette formule est valable à la fois pour les mots de code et pour les sous-parties d'arrêt. La condition $\omega > \omega_{\min}$ provient de la définition de l'ensemble expurgé. Afin de discuter les différents régimes résultant des diverses extrémisations, nous nous focalisons sur le cas des codes réguliers, qui permet un traitement analytique simplifié :

$$\begin{aligned} E_{BU} &= - \max_{e \in [0,1]} \left\{ -D(e \parallel \epsilon) + \min \left[\max_{\omega \in [\omega_{\min}, e]} \min_x (-\ell H(\omega) \right. \right. \\ &\quad \left. \left. - \omega \ell \log(x) + \frac{\ell}{k} \log p_{A,k}(x) + eH(\omega/e) \right), 0 \right] \right\} \end{aligned} \quad (5.60)$$

À mesure que le niveau de bruit ϵ augmente, le système passe par trois régimes de la phase décodable. Quand ϵ est faible, le maximum sur ω est atteint à la frontière $\omega_{\min}$. Comme dans les codes aléatoires, l'erreur y est dominée par une phase « condensée » où l'échec est presque toujours causé par un petit nombre de mots de code, ou de sous-parties d'arrêt, de poids $|S|$ minimal. L'événement rare causant l'erreur n'est pas tant un bruit élevé qu'un bruit compatible avec une A -partie de faible poids. Ce régime correspond au plancher d'erreur [MP03] que nous avons déjà évoqué. À l'autre extrême, pour ϵ relativement grand, la borne d'union explose et est remplacée par 1. Cette explosion signale l'existence d'un nombre élevé de A -parties,

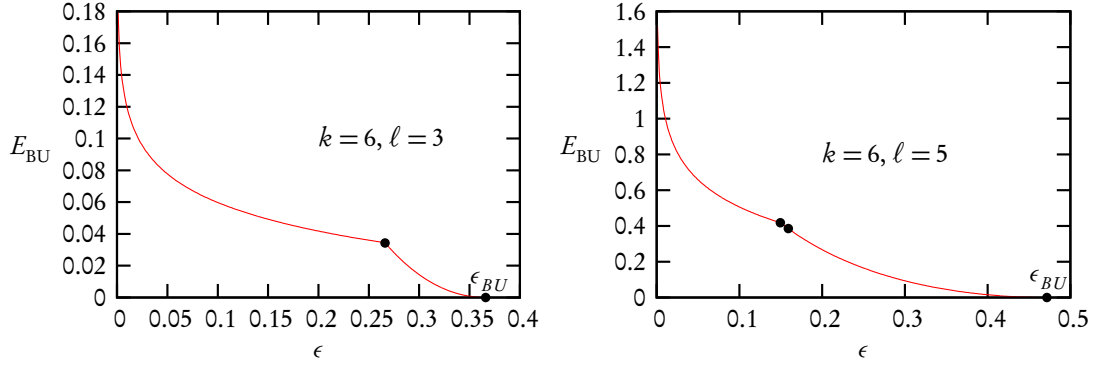


Fig. 5.6: Bornes d'union sur l'erreur du décodage itératif pour des codes réguliers avec $k=6, \ell=3$ (gauche) et $k=6, \ell=5$ (droite). Dans le premier cas le régime intermédiaire n'existe pas.

permises par un bruit élevé. Si on dénote par ϵ_{BU} le plus grand e tel que la somme sur $S \subset E$ n'explose pas, E_{BU} s'écrit dans ce régime $D(\epsilon_{BU}||\epsilon)$: l'erreur s'interprète alors comme provenant d'un bruit anormalement élevé : c'est la zone de « cascade ». Le seuil ϵ_{BU} est une borne inférieure au seuil réel ϵ_d (itératif, pour les sous-parties d'arrêt) ou ϵ_c (optimal, pour les mots de code). Enfin, dans une région intermédiaire des ϵ , l'extremum est atteint à l'intérieur du domaine (ω, e) , traduisant un équilibre entre la taille de l'A-partie fautive et le bruit.

La figure 5.6 représente la borne d'union sur l'erreur du décodage itératif, pour deux constructions de codes réguliers. Le cas du décodage optimal (cf. figure 5.7) est l'occasion de faire une comparaison avec l'exposant d'erreur du modèle à codes aléatoires (§1.2.4). Dans la limite $k, \ell \rightarrow \infty$, avec $k/\ell = 1 - R$, la borne d'union s'écrit :

$$E_{BU} = \begin{cases} -\delta_{GV}(R) \log \epsilon & \text{si } \epsilon < \frac{\delta_{GV}(R)}{1-\delta_{GV}} \\ 1-R - \log(1+\epsilon) & \text{si } \frac{\delta_{GV}(R)}{1-\delta_{GV}(R)} < \epsilon < \frac{1-R}{1+R}, \\ D(1-R||\epsilon) & \text{if } \frac{1-R}{1+R} < \epsilon < 1-R. \end{cases} \quad (5.61)$$

Cet exposant présente une différence notable avec celui du code aléatoire, cf. (1.87) : l'existence d'un régime de petit bruit, où l'erreur se condense sur les mots de code de distance minimale $\delta_{GV}(R)$. Cette condensation est l'effet propre de l'expurgation, qui améliore donc la performance du code. On conjecture même que l'ensemble aléatoire expurgé minimise la probabilité d'erreur sur tous les codes possibles, quand $N \rightarrow \infty$, saturant ainsi la fonction de fiabilité du canal.

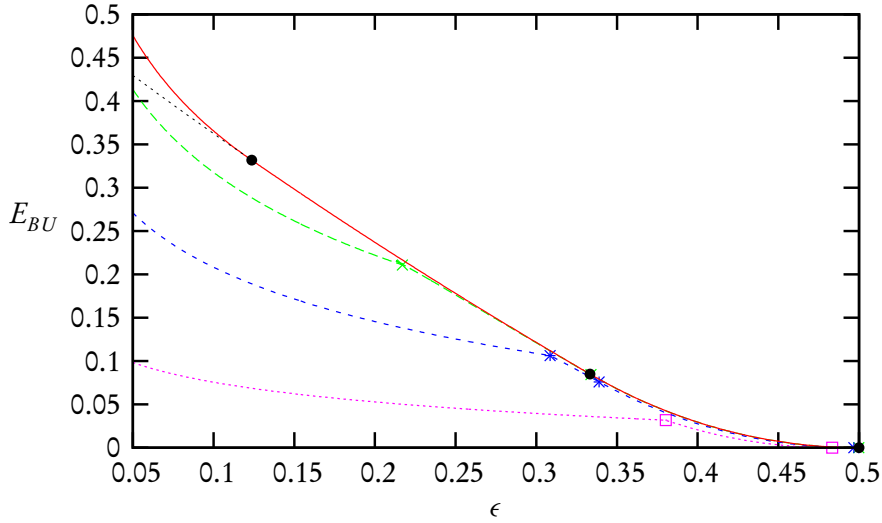


Fig. 5.7: Bornes d'union sur l'erreur du décodage optimal pour des codes réguliers de taux $R = 1/2$ et de degrés croissants : $(6, 3)$, $(8, 4)$, $(12, 6)$ et la limite des codes linéaires aléatoires (3.17). La ligne noire pointillée représente l'exposant de l'ensemble linéaire aléatoire non-expurgé.

Références

Le calcul recuit généralise les résultats de [DRU06] (mots de codes) et de [OVZ05] (sous-parties d'arrêt). L'introduction de l'ensemble lâche permet de faire le contact avec d'autres calculs recuits classiques, notamment dans l'ensemble poissonien.

La séparabilité des amas dans k -XORSAT a d'abord été prouvée dans [MRTZ03], où une erreur s'était glissée dans le calcul recuit de la fonction d'énumération des poids. Les valeurs de cette fonction, représentées figure 5.2, ont depuis été confrontées à des énumérations numériques exhaustives, produites par Federico Ricci-Tersenghi, avec un accord satisfaisant.

La preuve de la fragmentation dans k -SAT est reprise en grande partie de [MMZ05a, MMZ05b]. Le traitement du modèle jouet, qui n'a d'intérêt que pédagogique, a pour but de mettre en évidence les limites de l' x -satisfaisabilité. Notez bien que si les fonctions d'énumération des poids dans k -XORSAT, ou des distances dans k -SAT, se prêtent bien aux calculs du premier et du second moment, l'évaluation de leur moyenne gelée est beaucoup plus difficile. La fragmentation de l'espace des solutions impose en effet de recourir à une méthode de cavité avec une multiplicité d'états, comme nous allons le voir au chapitre suivant.

L'utilisation de la fonction d'énumération des poids pour les bornes d'union est classique en théorie de l'information [BM04]. Elle permet d'éclairer à l'aide d'outils simples la nature des différentes phases identifiées dans [MR06a].

Chapitre 6

Statistique des amas

Les chapitres précédents ont mis en évidence la nécessité de décrire, dans certains cas, l'espace de solutions comme une superposition d'états disjoints. Ce chapitre introduit une méthode permettant de concilier cette image, fondée sur la notion d'ordre à longue portée, avec un traitement local basé sur le passage des messages. La méthode est d'abord appliquée aux problèmes de satisfaction de contraintes, où le rôle des variables gelées est discuté. Elle sert ensuite à élucider les propriétés de distances dans les phases fragmentées.

6.1 Statistique des convictions

6.1.1 Une mesure sur les états

La discussion du chapitre 4 postule que la multiplicité des états équivaut à une multiplicité des solutions aux équations de Bethe¹ :

$$p_{i \rightarrow a}(\sigma_i) = \hat{p} \left[\{q_{b \rightarrow i}\}_{b \in i \setminus a} \right] (\sigma_i) \doteq 2^{F_{i \rightarrow a}} \prod_{b \in \partial i \setminus a} q_{b \rightarrow i}(\sigma_i), \quad (6.1)$$

$$q_{b \rightarrow i}(\sigma_i) = \hat{q} \left[\chi_b, \{p_{j \rightarrow b}\}_{j \in b \setminus i} \right] (\sigma_i) \doteq \sum_{\sigma_{b \setminus i}} \chi_b(\sigma_b) \prod_{j \in \partial b \setminus i} p_{j \rightarrow b}(\sigma_j). \quad (6.2)$$

À chaque état c est associée une solution des équations de cavité $\{p_{i \rightarrow a}^c, q_{a \rightarrow i}^c\}$, et vice-versa. Afin d'étudier la statistique des ces solutions, une mesure de Boltzmann est introduite :

$$P_m(c) = \frac{1}{\mathcal{Z}(m)} 2^{-m F_c}, \quad (6.3)$$

¹Par souci de clarté, nous nous permettrons par la suite d'omettre la température inverse β en la fixant à 1.

qui échantillonne les états selon leur énergie libre² F_c , ainsi que la fonction de potentiel associée

$$\phi(m) = \frac{1}{N} \log \mathcal{Z}(m) = \frac{1}{N} \log \sum_c 2^{-mF_c(\beta)}, \quad (6.4)$$

reliée à la complexité $\Sigma(f) = \frac{1}{N} \sum_c \mathbb{I}[F_c = Nf]$ par une transformée de Legendre :

$$\phi(m) = \frac{1}{N} \log \sum_f 2^{N[\Sigma(f) - m\beta f]} = \max_f [\Sigma(f) - m\beta f] \quad (6.5)$$

Transposée à l'espace des convictions, la mesure s'exprime comme :

$$P_m(\{p_{i \rightarrow a}, q_{a \rightarrow i}\}) = \frac{1}{\mathcal{Z}(m)} \prod_i 2^{-mF_{i+a \in \partial i}} \prod_a 2^{-mF_a} \prod_{(i,a)} 2^{mF_{ia}} \\ \times \prod_{(i,a)} \left\{ \delta[p_{i \rightarrow a} - \hat{p}(\{q_{b \rightarrow i}\}_{b \in \partial i \setminus a})] \delta[q_{a \rightarrow i} - \hat{q}(\chi_a, \{p_{j \rightarrow a}\}_{j \in \partial a \setminus i})] \right\} \quad (6.6)$$

$$\text{avec } 2^{-F_{i+a \in \partial i}} = \sum_{\sigma_i} \prod_{a \in \partial i} q_{a \rightarrow i}(\sigma_i), \quad 2^{-F_a} = \sum_{\sigma_a} \chi_a(\sigma_a) \prod_{i \in \partial a} p_{i \rightarrow a}(\sigma_i), \quad (6.7)$$

$$2^{-F_{ia}} = \sum_{\sigma_i} p_{i \rightarrow a}(\sigma_i) q_{a \rightarrow i}(\sigma_i) \quad (= 2^{-F_a} \text{ au point de cavité}) \quad (6.8)$$

Les facteurs « contraignant » de Dirac imposent que les équations de cavité soient vérifiées. S'y joignent des facteurs « souples » qui pondèrent ces solutions selon leur énergie libre, cf. (4.22).

6.1.2 Propagation des sondages

Cette mesure sur les convictions peut elle-même être étudiée au moyen de la méthode de la cavité. La figure 6.1 illustre comment s'organisent les différents facteurs de (6.6). Tout d'abord, les variables (p, q) sont disposées sur chaque lien (ia) : ce sont les « configurations » de notre problème³. À chaque nœud i on associe le jeu des contraintes $\delta(p_{i \rightarrow a} - \hat{p})$, pour $a \in \partial i$, ainsi que le facteur $2^{-mF_{i+a \in \partial i}}$. De même, on fait siéger les contraintes $\delta(q_{i \rightarrow a} - \hat{q})$, $i \in \partial a$, sur le nœud a , ainsi que le poids 2^{-mF_a} . Enfin, le poids $2^{-mF_{ia}}$, qui ne dépend que de la variable (p, q) , est disposé sur le nœud (ia) : il agit sur la variable comme un « champ extérieur⁴ ». Les équations brutes de cavité s'obtiennent comme la stricte application de l'approximation des arbres décrite au chapitre 4 :

²La quantité $-F_c$ est remplacée par l'entropie interne S_c dans la limite d'énergie nulle.

³Ces configurations sont ici de nature *continue*, à la différence des tous les cas discrets considérés jusqu'ici.

⁴C'est-à-dire un facteur de degré 1.

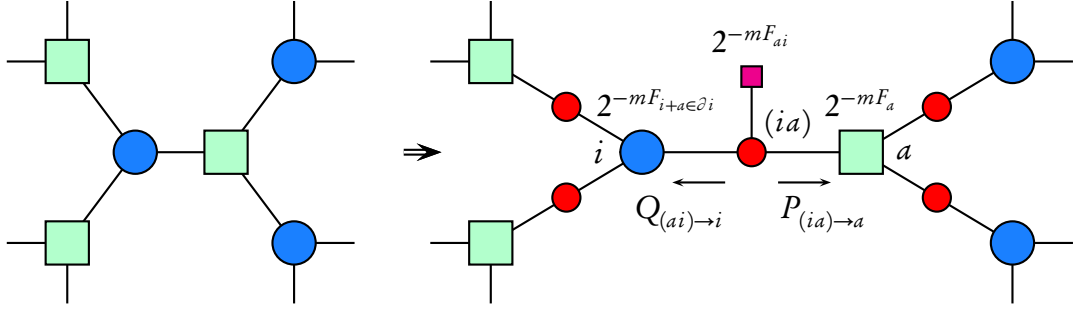


Fig. 6.1: Transformation du graphe factoriel initial en un graphe factoriel adapté à la mesure (6.6). Les nœuds i jouent maintenant le rôle de facteurs, au même titre que les nœuds a . Les variables sont sur les liens (ia) . Le schéma indique la localisation des poids d'énergie libre. Quant aux facteurs de Dirac, ils sont disposés sur les nœuds i , pour $\delta(p_{i \rightarrow a} - \hat{p})$, $a \in \partial i$, et sur les nœuds a , pour $\delta(q_{a \rightarrow i} - \hat{q})$, $i \in \partial a$.

$$\begin{aligned}
 P_{(ia) \rightarrow a}(p_{i \rightarrow a}, q_{a \rightarrow i}) &\propto \int \prod_{b \in \partial i \setminus a} [dp_{i \rightarrow b} dq_{b \rightarrow i} Q_{(bi) \rightarrow i}(p_{i \rightarrow b}, q_{b \rightarrow i})] \\
 &\times \delta[p_{i \rightarrow a} - \hat{p}(\{q_{b \rightarrow i}\}_{b \in \partial i \setminus a})] 2^{-m(F_{i+a \in \partial i} - F_{ia})} \\
 &\times \prod_{b \in \partial i \setminus a} \delta[p_{i \rightarrow b} - \hat{p}(q_{a \rightarrow i}, \{q_{c \rightarrow i}\}_{c \in \partial i \setminus a, b})]
 \end{aligned} \quad (6.9)$$

$$\begin{aligned}
 P_{(ia) \rightarrow a}(p_{i \rightarrow a}, q_{a \rightarrow i}) &\propto \int \prod_{b \in \partial i \setminus a} \{dq_{b \rightarrow i} Q_{(bi) \rightarrow i}[\hat{p}(q_{a \rightarrow i}, \{q_{c \rightarrow i}\}_{c \in \partial i \setminus a, b}), q_{b \rightarrow i}]\} \\
 &\times \delta[p_{i \rightarrow a} - \hat{p}(\{q_{b \rightarrow i}\}_{b \in \partial i \setminus a})] 2^{-mF_{i \rightarrow a}}
 \end{aligned} \quad (6.10)$$

$$\begin{aligned}
 Q_{(bi) \rightarrow i}(p_{i \rightarrow b}, q_{b \rightarrow i}) &\propto \int \prod_{j \in \partial b \setminus i} [dp_{j \rightarrow b} dq_{b \rightarrow j} P_{(jb) \rightarrow b}(p_{j \rightarrow b}, q_{b \rightarrow j})] \\
 &\times \delta[q_{b \rightarrow i} - \hat{q}(\chi_b, \{p_{j \rightarrow b}\}_{j \in \partial b \setminus i})] 2^{-m(F_b - F_{ib})} \\
 &\times \prod_{j \in \partial b \setminus i} \delta[q_{b \rightarrow j} - \hat{q}(\chi_b, p_{i \rightarrow b}, \{p_{k \rightarrow b}\}_{k \in \partial b \setminus i, j})]
 \end{aligned} \quad (6.11)$$

$$\begin{aligned}
 Q_{(bi) \rightarrow i}(p_{i \rightarrow b}, q_{b \rightarrow i}) &\propto \int \prod_{j \in \partial b \setminus i} \{dp_{j \rightarrow b} P_{(jb) \rightarrow b}[p_{j \rightarrow b}, \hat{q}(\chi_b, p_{i \rightarrow b}, \{p_{k \rightarrow b}\}_{k \in \partial b \setminus i, j})]\} \\
 &\times \delta[q_{b \rightarrow i} - \hat{q}(\chi_b, \{p_{j \rightarrow b}\}_{j \in \partial b \setminus i})]
 \end{aligned} \quad (6.12)$$

où on a utilisé le fait que $F_{i \rightarrow a}$, défini équation (4.17), vaut $F_{i+a \in \partial i} - F_a$, et que $F_{ai} = F_a$, au point fixe de la cavité.

Le recours à l'approximation des arbres peut paraître contradictoire, car l'hypothèse de la multiplicité des états a justement pour objectif de surmonter les difficultés liées à cette approximation. Si les formules de cavité sont bien exactes sur les arbres, les résultats qu'elles produisent peuvent en revanche être très sensibles aux conditions aux bords. Ainsi, les équations sur l'arbre infini peuvent admettre plusieurs solutions dont la statistique est auto-reproductrice. C'est précisément la statistique de ces solutions qu'explore la méthode de la cavité avec multiplicité d'états.

Une hypothèse supplémentaire permet de simplifier ces équations : la loi jointe $P_{(ia) \rightarrow a}$, renotée $P_{i \rightarrow a}$ par souci de concision, est supposée indépendante de $q_{a \rightarrow i}$. Cela revient à affirmer qu'en l'absence de a , aucune connaissance *a priori* sur le message $q_{a \rightarrow i}$ ne peut être extraite des messages provenant des autres voisins de i : en effet, dans la propagation des convictions, $q_{a \rightarrow i}$ se calcule à partir des messages provenant de la direction inverse. Notre hypothèse traduit ainsi la préservation de la causalité dans l'actualisation des convictions. L'hypothèse symétrique consiste naturellement à postuler que $Q_{(ai) \rightarrow i} \doteq Q_{a \rightarrow i}$ ne dépend que de $q_{a \rightarrow i}$. Ces deux hypothèses sont cohérentes entre elles, grâce au fait que $\Delta F_{i \rightarrow a}$ ne dépend que des $q_{b \rightarrow i}$.

Les équations précédentes prennent alors la forme classique [MP01] des équations de cavité avec multiplicité d'états⁵ :

$$P_{i \rightarrow a}(p_{i \rightarrow a}) \propto \int \prod_{b \in \partial i \setminus a} dq_{b \rightarrow i} Q_{b \rightarrow i}(q_{b \rightarrow i}) \delta \left[p_{i \rightarrow a} - \hat{p}(\{q_{b \rightarrow i}\}_{b \in \partial i \setminus a}) \right] 2^{-m F_{i \rightarrow a}} \quad (6.13)$$

$$Q_{b \rightarrow i}(q_{b \rightarrow i}) = \int \prod_{j \in \partial b \setminus i} dp_{j \rightarrow b} P_{j \rightarrow b}(p_{j \rightarrow b}) \delta \left[q_{b \rightarrow i} - \hat{q}(\chi_b, \{p_{j \rightarrow b}\}_{j \in \partial b \setminus i}) \right] \quad (6.14)$$

L'implémentation de ces équations comme règles d'actualisation donne lieu à une classe d'algorithmes désignés sous le terme anglais de *survey propagation* [MPZ02, MZ02], que nous traduisons par « propagation des sondages ».

La méthode de la cavité permet également d'évaluer la constante de normalisation $\mathcal{Z}(m)$, ou de manière équivalente la fonction de potentiel :

$$N \psi(m) = \sum_i \log \mathbb{E} \left(2^{-m F_{i+a \in \partial i}} \right) + \sum_a \log \mathbb{E} \left(2^{-m F_a} \right) - \sum_{(ia)} \log \mathbb{E} \left(2^{-m F_{ia}} \right) \quad (6.15)$$

⁵ « à un pas de brisure de symétrie des répliques » (*one-step replica symmetry breaking*).

avec :

$$\mathbb{E} \left(2^{-m F_{i+a \in \partial i}} \right) = \int \prod_{a \in \partial i} dq_{a \rightarrow i} Q_{a \rightarrow i}(q_{a \rightarrow i}) 2^{-m F_{i+a \in \partial i}}, \quad (6.16)$$

$$\mathbb{E} \left(2^{-m F_a} \right) = \int \prod_{i \in \partial a} dp_{i \rightarrow a} P_{i \rightarrow a}(p_{i \rightarrow a}) 2^{-m F_a}, \quad (6.17)$$

$$\mathbb{E} \left(2^{-m F_{ia}} \right) = \int dp_{i \rightarrow a} P_{i \rightarrow a}(p_{i \rightarrow a}) dq_{a \rightarrow i} Q_{a \rightarrow i}(q_{a \rightarrow i}) 2^{-m F_{ia}}. \quad (6.18)$$

Ces deux dernières quantités coïncident au point de fixe des équations de propagation des sondages.

Naturellement, en tant qu'équations de cavité sur la mesure (6.6), les équations avec multiplicité d'états possèdent toutes les propriétés déjà évoquées au chapitre 4, au premier rang desquelles leur caractère variationnel. Par ailleurs, on peut ici encore montrer [MP01], que la version moyennée de ces équations, qui font intervenir des distributions (sur les liens) de distributions (sur les états), sont équivalentes aux équations des répliques avec un pas de brisure de symétrie [Mon98].

6.1.3 Réduction à un état unique et condensation

Le choix de température interne $m = 1$ correspond en principe au calcul de l'énergie libre totale. En effet :

$$\psi(1) = \frac{1}{N} \log \sum_c 2^{-\beta F_c} = \frac{1}{N} \log \sum_c \sum_{\sigma \in c} 2^{-\beta E(\sigma)} = -\beta F(\beta) \quad (6.19)$$

Ce choix permet de réduire considérablement la complexité des équations avec multiplicité d'états. Définissant :

$$\bar{p}_{i \rightarrow a} = \int dp_{i \rightarrow a} P_{i \rightarrow a}(p_{i \rightarrow a}) p_{i \rightarrow a} \quad \bar{q}_{a \rightarrow i} = \int dq_{a \rightarrow i} Q_{a \rightarrow i}(q_{a \rightarrow i}) q_{a \rightarrow i} \quad (6.20)$$

il est facile de vérifier que ces moyennes vérifient précisément les équations de cavité à un état (4.17), (4.18). De la même façon, l'énergie libre totale est donnée par

$$\psi(1) = -\beta F(\{\bar{p}_{i \rightarrow a}, \bar{q}_{a \rightarrow i}\}) \quad (6.21)$$

où F est donnée par l'expression (4.29).

Cependant, ainsi que nous avons déjà eu l'occasion de le faire remarquer (§2.3.2), la température interne d'équilibre ne vaut pas nécessairement 1. Quand la complexité en $m = 1$,

$$\Sigma(m = 1) = -\partial_m \left(\frac{\psi(m)}{m} \right) \Big|_{m=1} \quad (6.22)$$

est négative, le système doit être décrit par une température interne plus élevée ($m^* < 1$), correspondant au seuil de condensation où la complexité s'annule. L'énergie libre réelle vaut alors :

$$-F = \frac{\psi(m^*)}{m^*}. \quad (6.23)$$

En revanche, dans la phase « liquide fragmentée » ($m = 1$), le système peut alternativement être décrit par une multiplicité d'états, ou par un état unique, ainsi que nous l'avons mentionné sans justification au paragraphe 2.3.2 pour k -SAT.

À cet égard, il est remarquable que le problème k -SAT connaisse, pour $\alpha \gtrsim 2^{k-1} \ln 2$, un régime où la mesure est fragmentée en une partition d'amas extensivement séparés, tout en restant, pour $\alpha \lesssim 2^k \ln 2 - (3/2) \ln 2$, appréhendable par un état unique.

6.1.4 Le seuil de satisfaisabilité

Rappelons que dans un problème de satisfactions de contraintes (où l'entropie remplace $-\beta F$), le seuil de satisfaisabilité est donné par l'annulation de la complexité totale $\psi(0) = \Sigma_{\text{tot}} = \max_s \Sigma(s)$, sous l'hypothèse de la fragmentation, cf. 2.3.2.

Un rapide coup d'œil à l'équation (6.15) pourrait laisser penser que le potentiel $\psi(m=0)$ s'annule toujours. En fait, il faut garder à l'esprit que les facteurs de pondération $2^{-F_{i+a \in \partial i}}$, 2^{-F_a} , etc., peuvent s'annuler du moment que les convictions peuvent elles-mêmes le faire pour certaines couleurs $\sigma \in \mathcal{X}$. La question de l'annulation ou non des convictions contient même toute l'information nécessaire, vu que les facteurs de pondération, pris à la puissance $m=0$, se comportent comme des fonctions de Dirac. Quand une conviction interdit certaines couleurs, on dit qu'elle contient un « avertissement ». L'annulation d'un facteur de pondération 2^s signale alors une contradiction entre des avertissements.

Plus formellement, les avertissements sont définis par $m_{i \rightarrow a} = \{\sigma | p_{i \rightarrow a}(\sigma) \neq 0\}$, et $n_{a \rightarrow i} = \{\sigma | q_{a \rightarrow i}(\sigma) \neq 0\}$, et les contraintes de propagation s'écrivent :

$$m_{i \rightarrow a} = \bigcap_{b \in \partial i \setminus a} n_{b \rightarrow i} \quad (6.24)$$

$$n_{a \rightarrow i} = \left\{ \sigma_i \mid \exists \sigma_{a \setminus i}, \sigma_j \in m_{j \rightarrow a}, \text{ t. q. } \chi(\sigma_a) \neq 0 \right\} \quad (6.25)$$

S'y ajoute une clause de non-contradiction :

$$m_i \doteq \bigcap_{a \in \partial i} n_{a \rightarrow i} \neq \emptyset, \quad \text{i.e.} \quad n_{a \rightarrow i} \cap m_{i \rightarrow a} \neq \emptyset. \quad (6.26)$$

Si m_i est un singleton, alors la variable i est gelée. De même, si $m_{i \rightarrow a}$ est un singleton, on dira que i est gelée en l'absence de a .

Dans le cas binaire $\sigma_i \in \{0, 1\}$, il n'existe que trois cas possibles pour $m_{i \rightarrow a}$: interdiction de 1, interdiction de 0, ou aucune interdiction. Les probabilités de ces trois événements sont respectivement notées $\pi_{i \rightarrow a}^0$, $\pi_{i \rightarrow a}^1$ et $\pi_{i \rightarrow a}^*$:

$$P_{i \rightarrow a}(p_{i \rightarrow a}) = \pi_{i \rightarrow a}^0 \delta(p_{i \rightarrow a}, \delta_0) + \pi_{i \rightarrow a}^1 \delta(p_{i \rightarrow a}, \delta_1) + \pi_{i \rightarrow a}^* \tilde{P}_{i \rightarrow a} \quad (6.27)$$

où $\tilde{P}_{i \rightarrow a}$ est une loi de support $]0, 1[$. À $m = 0$, l'algèbre des avertissements se suffit à elle-même : la partie « indécise⁶ » des convictions, représentée par $\tilde{P}_{i \rightarrow a}$, contient des détails spécifiques à l'amas considéré, mais n'influence pas la statistique uniforme qui ignore la structure interne des amas.

Par souci de lisibilité, on utilisera par la suite les abbréviations suivantes : $\{0, 1\} = *$, $\{0\} = 0$ et $\{1\} = 1$.

k -XORSAT

L'algèbre des avertissements prend une forme simple dans le cas des équations booléennes linéaires. Elle caractérise en fait le point fixe de l'algorithme d'effeuillage décrit au §3.2.1, comme nous allons le voir.

Notons $\eta_{a \rightarrow i}^0$, $\eta_{a \rightarrow i}^1$ et $\eta_{a \rightarrow i}^*$ les probabilités que $n_{a \rightarrow i}$ vaille respectivement 0, 1 ou * en l'absence de i . On a :

$$\eta_{a \rightarrow i}^\sigma = \frac{1}{2} \left[\prod_{j \in \partial a \setminus i} (\pi_{j \rightarrow a}^0 + \pi_{j \rightarrow a}^1) + (-1)^{\sigma + \tau_a} \prod_{j \in \partial a \setminus i} (\pi_{j \rightarrow a}^0 - \pi_{j \rightarrow a}^1) \right] \quad (6.28)$$

Pour ce qui est du deuxième type des équations de propagation des avertissements, la variable i envoie l'avertissement $m_{i \rightarrow a} = \sigma$ en l'absence de a si au moins l'un de ses autres tests lui commande de valoir σ , et aucun ne lui demande de valoir $1 - \sigma$:

$$\pi_{i \rightarrow a}^\sigma \propto \prod_{b \in \partial i \setminus a} (1 - \eta_{b \rightarrow i}^{1-\sigma}) - \prod_{b \in \partial i \setminus a} \eta_{b \rightarrow i}^* \quad (6.29)$$

$$\pi_{i \rightarrow a}^* \propto \prod_{b \in \partial i \setminus a} \eta_{b \rightarrow i}^* \quad (6.30)$$

La construction des amas par l'algorithme d'effeuillage entraîne que les variables gelées doivent être les mêmes *pour tous les amas*. Ce postulat est-il compatible avec les équations de propagation des sondages données ci-dessus ?

Commençons avec des conditions initiales symétriques $\pi_{i \rightarrow a}^0 = \pi_{i \rightarrow a}^1 = 1/2$, et $\eta_{a \rightarrow i}^0 = \eta_{a \rightarrow i}^1 = 1/2$, et itérons la propagation des sondages. Au premier pas, seuls

⁶Dans le langage de la physique statistique, une conviction indécise correspond à un « champ évanescent », tandis qu'un avertissement symbolise un « champ dur ».

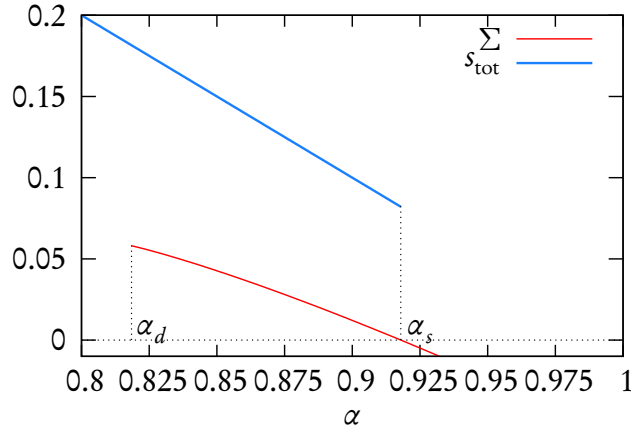


Fig. 6.3: Complexité et entropie totale de 3-XORSAT en fonction de la densité de tests $\alpha = M/N$.

de point fixe de l'algorithme BP pour un code linéaire dilué sous le canal d'effacement (4.46), (4.47), qui décrivent également un algorithme d'effeuillage.

Ce raisonnement, en ce qu'il repose sur des arguments purement géométriques, ne dépend pas de m . Il a déjà été argumenté que l'entropie interne ne doit pas non plus dépendre de l'amas considéré. Vérifions que c'est bien le cas ici. Le même type d'argument que précédemment conduit à supposer la symétrie entre 0 et 1 au sein d'un même état, c'est-à-dire dans la partie indécidée des convictions. Formellement cette condition s'écrit :

$$\tilde{P}_{i \rightarrow a} = \delta_{1/2}. \quad (6.34)$$

Ainsi, toute conviction qui n'est pas un avertissement est parfaitement équilibrée. Cette hypothèse, qui est cohérente avec elle-même, est également compatible avec ce que l'on sait du processus de reconstruction, qui n'opère que des choix symétriques.

L'injection de cet Ansatz dans la version moyennée de (6.15) permet de calculer la fonction de potentiel dans la limite $N \rightarrow \infty$:

$$\psi(m) = 1 - (1 + \lambda)e^{-\lambda} - \alpha(1 - e^{-\lambda})^k + m \left[(1 + \lambda)e^{-\lambda} + \alpha \left((1 - e^{-\lambda})^k - 1 \right) \right] \quad (6.35)$$

Comme attendu, on trouve que l'entropie interne, donnée par $\partial_m \psi(m)$, ne dépend pas de m . L'entropie totale vaut $\psi(1) = 1 - \alpha$, ainsi que l'avons vu au paragraphe 3.2.1. L'annulation de la complexité $\Sigma = \psi(0)$ donne le seuil de satisfaisabilité α_s , cf. figure 6.3.

k -SAT

Dans k -SAT, l'avertissement $n_{a \rightarrow i}$ ne peut prendre que les valeurs $*$ ou σ_i^a , où σ_i^a désigne le littéral de i dans a . Autrement dit, une clause a peut commander à i de la

satisfaire, ou ne rien lui demander. La clause a ne contraindra i que si *chaque* autre voisin de a lui envoie un avertissement la prévenant qu'il ne pourra pas la satisfaire. Cela se produit avec probabilité :

$$\eta_{a \rightarrow i} = \prod_{j \in \partial a \setminus i} \pi_{j \rightarrow a}^n \quad (6.36)$$

où $\pi_{j \rightarrow a}^n = \pi_{j \rightarrow a}^{\bar{\sigma}^j}$. Par ailleurs, une variable i est gelée (en l'absence de a) en σ si et seulement si :

- Clause de gel : au moins une des clauses où i apparaît avec le même littéral que σ (ce qu'on note $b \in \partial_{\sigma} i \setminus a$), envoie un avertissement.
- Clause de non-contradiction : aucune des clauses où i apparaît sous la forme du littéral inverse de σ ($b \in \partial_{\bar{\sigma}} i \setminus a$), n'envoie d'avertissement.

Parallèlement, la variable i sera non-gelée (en l'absence de a) si aucune des autres clauses de i ($b \in \partial i \setminus a$) n'envoie d'avertissement. Pour résumer :

$$\pi_{i \rightarrow a}^{\sigma} \propto \left[1 - \prod_{b \in \partial_{\sigma} i \setminus a} (1 - \eta_{b \rightarrow i}) \right] \prod_{b \in \partial_{\bar{\sigma}} i \setminus a} (1 - \eta_{b \rightarrow i}) \quad (6.37)$$

$$\pi_{i \rightarrow a}^* \propto \prod_{b \in \partial i \setminus a} (1 - \eta_{b \rightarrow i}) \quad (6.38)$$

Ces équations sont exactement les équations de propagation des sondages (*survey propagation*, [BMZ05]) telles qu'elles ont été introduites pour la première fois dans le problème k -SAT.

La complexité $\Sigma_{\text{tot}} = \psi(m=0)$ vaut quant à elle :

$$\Sigma_{\text{tot}} = \sum_i \log(\Pi_i^0 + \Pi_i^1 + \Pi_i^*) - \sum_a (|\partial a| - 1) \log \left(1 - \prod_{i \in a} \pi_{i \rightarrow a}^n \right) \quad (6.39)$$

$$\text{où } \Pi_i^{\sigma} = \left[1 - \prod_{a \in \partial_{\sigma} i} (1 - \eta_{a \rightarrow i}) \right] \prod_{a \in \partial_{\bar{\sigma}} i} (1 - \eta_{a \rightarrow i}) \quad (6.40)$$

$$\Pi_i^* = \prod_{a \in \partial i} (1 - \eta_{a \rightarrow i}) \quad (6.41)$$

La moyennation de ces équations sur un graphe poissonien permet de calculer le seuil de satisfaisabilité $\alpha_s(k)$, comme le montre la figure 6.4.

Le traitement à m quelconque fait intervenir des objets nettement plus compliqués, car il faut inclure les détails des convictions indéfinies, dont la distribution est

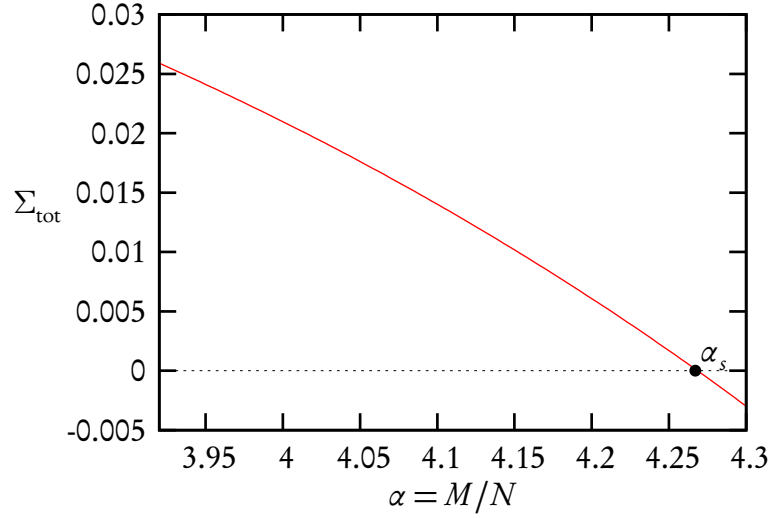


Fig. 6.4: Complexité totale de 3-SAT aléatoire en fonction de la densité de clauses $\alpha = M/N$.

donnée par $\tilde{P}_{i \rightarrow a}$. Néanmoins, dans la limite des grands k , il est possible d'effectuer le calcul de manière analytique. Dans k -SAT, les équations de cavité à un état s'expriment comme :

$$p_{i \rightarrow a}(\sigma_i) = 2^{-S_{i \rightarrow a}} \prod_{b \in \partial_{\sigma_i} i \setminus a} q_{b \rightarrow i}(\sigma_i) \quad (6.42)$$

$$q_{b \rightarrow i} \doteq q_{b \rightarrow i}(\bar{\sigma}_i^b) = 1 - \prod_{j \in \partial b \setminus i} p_{j \rightarrow b}(\bar{\sigma}_i^b), \quad q_{b \rightarrow i}(\sigma_i^b) = 1 \quad (6.43)$$

À grand k , on trouve que le point fixe de ces équations sous la mesure $2^{m S_{i \rightarrow a}}$ vérifie :

$$p_{i \rightarrow a} = \begin{cases} \delta_0 & \text{prob.} \sim 1/2, \\ \delta_1 & \text{prob.} \sim 1/2, \\ \frac{\delta_0 + \delta_1}{2} & \text{prob.} \sim 2^{-k-1} 2^m \end{cases} \quad q_{b \rightarrow i} = \begin{cases} 1 & \text{prob.} \sim 1, \\ 0 & \text{prob.} \sim 2^{1-k}, \\ \frac{1}{2} & \text{prob.} \sim (k-1) 2^{1-2k} 2^m \end{cases} \quad (6.44)$$

L'injection de cette solution dans l'expression du potentiel donne :

$$2^{k+1} \ln(2) \phi(m) = 2^m + 2[2^k \ln(2) - \alpha] - 2 - \ln(2) \quad (6.45)$$

On retrouve ainsi, par l'annulation de $\Sigma(m=1)$ et $\Sigma(m=0)$, le comportement asymptotique de $\alpha_c(k)$ et $\alpha_s(k)$ donné par le tableau 2.1. En passant, il est intéressant de noter qu'en posant

$$p = 1 - 2^{-k-1} \quad \text{et} \quad \alpha_{\text{sat}} = 2^k \ln(2) - \frac{1 + \ln(2)}{2} - 2^k \ln(2)(1 - \alpha_{\text{jouet}}),$$

on retrouve les expressions du modèle à amas aléatoires dans la limite des très petits amas. Ainsi, k -SAT «tend» vers un modèle à mots de code aléatoires ($p = 1$), la première correction asymptotique étant donnée par le modèle à amas aléatoires.

En particulier, il en découle que les variables libres sont asymptotiquement indépendantes entre elles, car $s \approx \pi_i^* \approx 2^m 2^{-k-1}$. Cette observation permet d'expliquer la bonne précision de la borne de premier moment décrite dans [DB97, KKKS98], qui donne $\alpha_s \lesssim 2^k \ln(2) - [1 + \ln(2)]/2$. Cette méthode se fonde sur le comptage des solutions «négativement premières», pour lesquelles aucune des variables $\sigma_i = 1$ ne peut être changée en 0 sans violer la formule. Il est facile de voir que dans la limite $k \rightarrow \infty$, chaque amas contient exactement une solution négativement première. La méthode du premier moment compte donc approximativement le nombre d'amas, au lieu de compter le nombre de solutions comme dans la borne naïve, cf. (2.15).

6.2 Modèles étendus

Dans les problèmes de satisfaction de contrainte, le seuil de satisfaisabilité découle uniquement de la statistique des avertissements. Ces derniers sont liés par un jeu de contraintes spécifiques induites par les contraintes du problème original.

Le nouveau problème de satisfaction de contraintes ainsi défini, appelé *modèle étendu*, retrouve une nature discrète, puisque ses variables, les avertissements (n, m) , se contentent d'indiquer des ensembles des couleurs interdites. Ces nouvelles variables doivent satisfaire les clauses suivantes :

$$\begin{aligned} m_{i \rightarrow a} &= \bigcap_{b \in \partial i \setminus a} n_{b \rightarrow i}, & n_{a \rightarrow i} \cap m_{i \rightarrow a} &\neq \emptyset, \\ n_{a \rightarrow i} &= \left\{ \sigma_i \mid \exists \sigma_{a \setminus i}, \sigma_j \in m_{j \rightarrow a}, \text{ t. q. } \chi(\sigma_a) \neq 0 \right\}. \end{aligned} \tag{6.46}$$

Les solutions du modèle étendu s'identifient grossièrement aux amas du problème original⁷, bien que la correspondance ne soit pas parfaite : par exemple la configuration sans avertissement $n_{a \rightarrow i} = \{1, \dots, q\}$, $m_{i \rightarrow a} = \{1, \dots, q\}$ est toujours solution du problème étendu, sans pour autant nécessairement correspondre à un amas, notamment dans la phase insatisfaisable. Par ailleurs, même dans l'hypothèse où un tel amas existerait, il n'est pas garanti qu'il soit unique.

Malgré cela, il peut être intéressant d'étudier le modèle étendu pour soi. On peut par exemple envisager de se servir des solutions du modèle étendu comme point de départ à la recherche de solutions au problème original, divisant ainsi la recherche de solutions en deux étapes distinctes. Une telle procédure permet d'exercer un contrôle

⁷C'est-à-dire les amas dominant la mesure uniforme $m = 0$.

sur la localisation de la solution, ainsi que sur son environnement. La caractérisation des amas pourrait également servir à concevoir de nouveaux procédés de correction d'erreur ou de compression [BBCZ05].

6.2.1 Fonction d'énumération du gel

L'énumération des solutions du modèle étendu redonne naturellement, dans le cadre de l'approximation de Bethe, la complexité totale $\psi(0)$. Il peut être intéressant d'explorer les grandes déviations de ces solutions, en variant notamment le nombre de variables qui y sont gelées. Notons $\mathcal{S}$ l'ensemble des solutions d'un problème étendu et, pour une solution (m, n) , désignons par $w(m, n)$ le nombre de variables i gelées, i.e. telles que m_i est un singleton. Nous introduisons la *fonction d'énumération de gel* :

$$Z(\beta) = \sum_{(m,n) \in \mathcal{S}} 2^{-\beta w(m,n)} \quad (6.47)$$

qui échantillonne les solutions du modèle étendu selon le nombre de variables gelées. Elle génère le nombre n_w d'amas ayant un nombre fixé w de variables gelées.

L'introduction du poids $2^{-\beta w}$ brise l'hypothèse d'indépendance qui avaient permis la simplification des équations de cavité au §6.1.2. En effet, l'introduction d'une pondération dépendant du gel est incompatible avec la directionnalité causale de la propagation des convictions : quand i envoie un message d'avertissement à a , il doit maintenant tenir compte de l'influence que a lui-même exerce sur son gel éventuel, déterminé par $m_i = n_{a \rightarrow i} \cap m_{i \rightarrow a}$. La variable $m_{a \rightarrow i}$ ne peut donc être ignorée, comme c'était le cas pour $\beta = 0$.

Rappelons que dans k -SAT, $n_{a \rightarrow i}$ ne peut prendre que deux valeurs : soit $*$, soit σ_i^a . Nous nous intéressons plus particulièrement aux quatre cas de figure suivants pour la variable $(m_{i \rightarrow a}, n_{a \rightarrow i})$, en l'absence de a :

- (σ_i^a, σ_i^a) ou $(*, \sigma_i^a)$: la variable i est gelée, entre autre par a qui lui envoie un avertissement.
- $(\sigma, *)$: la variable i est gelée en σ , mais a ne lui envoie pas d'avertissement.
- $(*, *)$: la variable i est libre.

Les probabilités de ces événements sous $P_{(ia) \rightarrow a}$ sont respectivement dénotées $\pi_{i \rightarrow a}^g$, $\pi_{i \rightarrow a}^\sigma$, et $\pi_{i \rightarrow a}^*$. Pour ce qui est des messages allant dans l'autre sens, quatre situations se dégagent pour $(m_{i \rightarrow a}, n_{a \rightarrow i})$:

$$\begin{aligned} \eta_{a \rightarrow i}^* &= Q_{(ai) \rightarrow i}(*, *), & \eta_{a \rightarrow i}^n &= Q_{(ai) \rightarrow i}(\bar{\sigma}_i^a, *) \\ \eta_{a \rightarrow i}^s &= Q_{(ai) \rightarrow i}(\sigma_i^a, *), & \eta_{a \rightarrow i}^g &= Q_{(ai) \rightarrow i}(*, \sigma_i^a) + Q_{(ai) \rightarrow i}(\sigma_i^a, \sigma_i^a) \end{aligned} \quad (6.48)$$

Ces jeux de messages sont reliés par les relations suivantes :

$$\pi_{i \rightarrow a}^* \propto 2^\beta \prod_{b \in \partial i \setminus a} \eta_{b \rightarrow i}^* \quad (6.49)$$

$$\pi_{i \rightarrow a}^\sigma \propto \prod_{b \in \partial_\sigma i \setminus a} \eta_{b \rightarrow i}^n \left[\prod_{b \in \partial_\sigma i \setminus a} (\eta_{b \rightarrow i}^s + \eta_{b \rightarrow i}^g) - \prod_{b \in \partial_\sigma i \setminus a} \eta_{b \rightarrow i}^s \right] \quad (6.50)$$

$$\pi_{i \rightarrow a}^g \propto \prod_{b \in \partial_{\sigma^a} i \setminus a} \eta_{b \rightarrow i}^n \prod_{b \in \partial_{\sigma^a} i \setminus a} (\eta_{b \rightarrow i}^s + \eta_{b \rightarrow i}^g) \quad (6.51)$$

Dans la deuxième équation, on impose qu'au moins un avertissement soit reçu par un voisin de i distinct de a . Cette condition n'est pas nécessaire dans la troisième équation, puisque a suffit à geler i .

L'autre volet des équations s'écrit, avec les notations $\pi_{j \rightarrow a}^s = \pi_{j \rightarrow a}^{\sigma_j^a}$, $\pi_{j \rightarrow a}^n = \pi_{j \rightarrow a}^{\bar{\sigma}_j^a}$:

$$\eta_{a \rightarrow i}^g = \prod_{j \in \partial a \setminus i} \pi_{j \rightarrow a}^n \quad (6.52)$$

$$\eta_{a \rightarrow i}^s = \eta_{a \rightarrow i}^* = \prod_{j \in \partial a \setminus i} (\pi_{j \rightarrow a}^0 + \pi_{j \rightarrow a}^1 + \pi_{j \rightarrow a}^*) - \prod_{j \in \partial a \setminus i} \pi_{j \rightarrow a}^n \quad (6.53)$$

$$\begin{aligned} \eta_{a \rightarrow i}^n &= \prod_{j \in \partial a \setminus i} (\pi_{j \rightarrow a}^0 + \pi_{j \rightarrow a}^1 + \pi_{j \rightarrow a}^*) - \prod_{j \in \partial a \setminus i} \pi_{j \rightarrow a}^n \\ &+ \sum_{j \in \partial a \setminus i} (\pi_{j \rightarrow a}^g - \pi_{j \rightarrow a}^* - \pi_{j \rightarrow a}^s) \prod_{j' \neq j} \pi_{j' \rightarrow a}^n \end{aligned} \quad (6.54)$$

Dans la dernière équation, on a pris soin de ne pas oublier que quand i envoie un avertissement contrariant la clause a , celle-ci peut en retour geler l'un de ses parents j .

Dans le cas particulier $\beta = 0$, l'hypothèse d'indépendance redevient valide. D'un côté, elle implique que la probabilité $P_{(ia) \rightarrow a}(m_{i \rightarrow a} \neq \bar{\sigma}_i^a)$ ne dépend pas de $n_{a \rightarrow i}$, ce qui se traduit par $\pi^g = \pi^* + \pi^s$. De l'autre côté, la probabilité $Q_{(ai) \rightarrow i}(n_{a \rightarrow i} = *)$ ne doit pas dépendre de $m_{i \rightarrow a}$, autrement dit : $\eta^* = \eta^n = \eta^s$. Ces deux conditions sont cohérentes tant que $\beta = 0$. On retrouve ainsi les équations de propagation des sondages (6.36)–(6.38).

La transformation de Legendre inverse de la fonction de potentiel $\phi(\beta) = \log Z(\beta)$ permet de remonter au taux de croissance de n_w , déterminé par l'entropie $\Omega(\omega) = \frac{1}{N} \log n_{N\omega}$. Ce taux est représenté figure 6.5 pour $k = 3$. Un petit calcul analytique montre que, dans la limite $k \rightarrow \infty$, il coïncide avec la complexité $\Sigma(s)$ déduite de (6.35) :

$$\Omega(\omega) \sim \Sigma(s = 1 - \omega), \quad \text{i.e.} \quad \phi(\beta) = \psi(m = \beta) - \beta \quad (6.55)$$

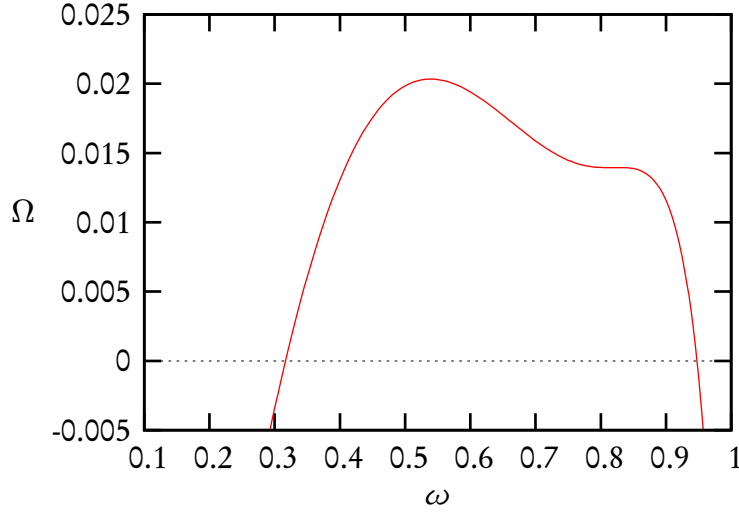


Fig. 6.5: Entropie ($n_w = 2^{N\Omega}$) des solutions du modèle étendu, en fonction du nombre de variables gelées $N\omega$, pour une formule aléatoire avec $k = 3$, $N = 10\,000$ et $M = 41\,000$. Le maximum de droite correspond à Σ_{tot} , qui gouverne le nombre total d'amas dans le problème original. Le maximum de gauche n'est pas physique, et est situé dans une région instable des équations de cavité.

Cette identité n'est guère surprenante : dans cette limite, le nombre de variables libres d'un amas égale son entropie $s = 2^m 2^{-k-1} = 2^\beta 2^{-k-1}$: les variables qui sont libres les sont indépendamment les unes des autres.

Il est intéressant de remarquer qu'une généralisation similaire des équations de propagation des sondages a déjà été proposée dans [MMW05], avec des poids plus généraux. Celle-ci donne lieu à une amélioration des performances de la procédure de décimation (*Survey Inspired Decimation*, [MZ02]).

6.2.2 Blanchissement

Une question naturelle concernant le modèle étendu est de savoir comment construire ses solutions de manière explicite. Si l'on dispose d'une solution σ^0 au problème original, il est facile de construire la solution du modèle étendu correspondant à l'amas contenant σ^0 . On implémente pour cela l'algorithme du *blanchissement*, qui n'est rien d'autre que l'algorithme de propagation des avertissements avec pour conditions initiales :

$$m_{i \rightarrow a} = \{\sigma_i^0\}. \quad (6.56)$$

On commence donc avec des messages colorés. Puis, à chaque pas de l'algorithme, les variables sous-contraintes (c'est-à-dire pour lesquelles plusieurs couleurs sont per-

mises) sont peu à peu « blanchies ». Ce blanchissement se propage jusqu'à atteindre un point fixe, qui ne dépend pas de l'ordre des actualisations. De toute évidence, l'amas décrit par ce point fixe contient la solution originale σ^0 . On note $w(\sigma^0)$ le nombre de variables gelées une fois le point fixe atteint.

Malheureusement, l'implémentation pratique du blanchissement dans k -SAT, à partir de solutions σ^0 issues d'algorithmes classiques (*Random WalkSAT*, *Survey Inspired Decimation*), mène presque toujours à la solution triviale $m = n = *$ pour les grands problèmes [MMW05].

L'analyse rigoureuse de l'algorithme de blanchissement a cependant permis de montrer le résultat suivant [ART06] : pour tout $\omega < 1$, et pour tout $k \geq k_0(\omega)$, il existe un seuil $\alpha_g(\omega, k) < \alpha_s(k)$ tel que pour tout $\alpha \in [\alpha_g(\omega, k), \alpha_s(k)]$, le nombre de solutions σ^0 telles que $w(\sigma^0) < N\omega$ s'annule presque sûrement. Autrement dit, tous les amas ont au moins ωN variables gelées, du moment que k est suffisamment grand — en l'occurrence, $k \geq 9$.

Malgré ce résultat positif, la question de l'existence d'une solution non-triviale au modèle étendu reste encore largement ouverte en toute généralité. Parmi les problèmes intéressants liés à cette question, la preuve de l'existence d'une transition abrupte, qui coïnciderait avec le seuil α_s du problème original, permettrait de valider le scénario de la multiplicité des états, et fournirait ainsi un socle solide à ses prédictions quantitatives.

6.3 Retour sur les distances

Dans les calculs effectués au chapitre 5, nous nous sommes la plupart du temps restreints à un traitement recuit des propriétés de distance. En effet, en présence d'une phase fragmentée, le calcul de la moyenne gelée doit recourir au formalisme de cavité avec multiplicité d'état. Nous mettons ici en œuvre ce formalisme, en tirant profit du cadre général présenté au §6.1.

6.3.1 Diamètre

Afin de déterminer le spectre des distances d'un problème de satisfactions de contraintes, il nous faut tout d'abord connaître la distribution des diamètres d'amas. On définit une fonction d'énumération interne des distances, pour un amas c donné :

$$Z(c, \beta) \doteq 2^{-\beta F(c, \beta)} = \sum_{(\sigma, \sigma') \in c^2} 2^{-\beta \|\sigma - \sigma'\|}, \quad (6.57)$$

La limite $\beta \rightarrow -\infty$ est dominée par la paire la plus éloignée de c . Autrement dit, elle donne son diamètre :

$$\lim_{\beta \rightarrow -\infty} F(c, \beta) = w_d(c), \quad \text{avec} \quad w_d(c) \doteq \max_{(\sigma, \sigma') \in c^2} \|\sigma - \sigma'\|. \quad (6.58)$$

La mesure adéquate permettant d'échantillonner les amas selon leur fonction d'énumération est définie par :

$$\mathcal{Z}_d(m, \beta) = \sum_c 2^{-m\beta F(c, \beta)}. \quad (6.59)$$

La limite $\beta \rightarrow -\infty$, avec $y \doteq m\beta$ constant, transforme cette mesure en un échantillonnage des diamètres :

$$\mathcal{Z}_d(y) \doteq 2^{\phi_d(y)} = \sum_c 2^{-y w_d(c)} \quad (6.60)$$

Le problème d'optimisation associé à la recherche de la paire la plus éloignée au sein d'un même amas est malheureusement assez mal défini du point de vue de la cavité. Comment, en effet, s'assurer que les solutions échantionnées appartiennent bien au même amas ? Une manière d'imposer cette condition consiste à forcer les deux solutions à suivre les prescriptions du même point fixe de la propagation des avertissements, autrement dit, à coïncider sur les variables gelées de la même solution du modèle étendu. Bien entendu, la fiabilité de cette recette dépend de la confiance qu'on accorde au modèle étendu pour représenter les amas du modèle original.

On suppose que dans un même amas, la mesure $Z(c, \beta)$ peut être décrite par le formalisme de cavité à un état :

$$p_{i \rightarrow a}(\sigma_i, \sigma'_i) \propto 2^{-\beta \delta_{\sigma_i, \sigma'_i}} \prod_{b \in \partial i \setminus a} q_{b \rightarrow i}(\sigma_i, \sigma'_i) \quad (6.61)$$

$$q_{b \rightarrow i}(\sigma_i, \sigma'_i) = \sum_{\sigma_{b \setminus i}, \sigma'_{b \setminus i}} \prod_{j \in \partial b \setminus i} p_{j \rightarrow b}(\sigma_j, \sigma'_j) \chi_b(\sigma_b) \chi_b(\sigma'_b) \quad (6.62)$$

Les contributions $w_{i+a \in \partial i}^d$, w_a^d et w_{ia}^d au diamètre $w_d(c)$ s'obtiennent par des équations semblables, dans la limite $\beta \rightarrow -\infty$. La condition sur les avertissements implique par ailleurs :

$$m_{i \rightarrow a} = \{\sigma_i \mid \sum_{\sigma'_i} p_{i \rightarrow a}(\sigma_i, \sigma'_i) \neq 0\} = \{\sigma'_i \mid \sum_{\sigma_i} p_{i \rightarrow a}(\sigma_i, \sigma'_i) \neq 0\} \quad (6.63)$$

Si l'on désigne par $\mathcal{A}_\beta(p, q)$ l'événement par lequel les conditions (6.61), (6.62) et (6.63) sont vérifiées, le potentiel (6.60) se réécrit :

$$\phi_d(y) = \frac{1}{N} \log \int dp \, dq \, \mathbb{I}[\mathcal{A}_{\beta \rightarrow -\infty}(p, q)] \prod_i 2^{-y w_{i+a \in \partial i}^d} \prod_a 2^{-y w_a^d} \prod_{(ia)} 2^{y w_{ia}^d} \quad (6.64)$$

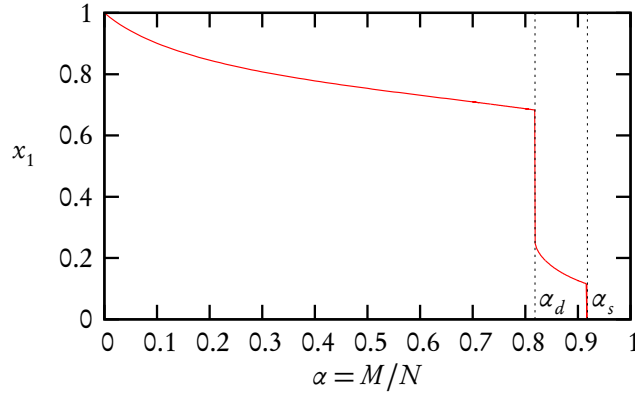


Fig. 6.6: Diamètre des amas du problème 3-XORSAT aléatoire en fonction de la densité de tests. En α_d , l'amas géant se fragmente en un grand nombre de petits amas, faisant subir une discontinuité au diamètre $x_1 = w_d/N$.

En principe, le diamètre maximal devrait être donné par la limite $y \rightarrow -\infty$. En pratique cependant, la complexité :

$$\Sigma_d(\omega) = \frac{1}{N} \log \sum_c \mathbb{I}(w_M(c, c) = N\omega) \quad (6.65)$$

reliée au potentiel $\phi_d(y)$ par une transformée de Legendre, s'annule à une température inverse y finie. Le modèle jouet d'amas aléatoires, qui a déjà été étudié au chapitre précédent, permet de s'en convaincre. La statistique des diamètres y suit une loi binomiale de paramètre $1 - p$. Ainsi on a (cf. fig. 6.7a) :

$$\Sigma_d(\omega) = 1 - \alpha - D(\omega || 1 - p). \quad (6.66)$$

Par Chebychev, le nombre d'amas est typiquement nul dès que la complexité est négative. La thermodynamique sera donc dominée par une température intermédiaire $y^* > -\infty$ signalant l'annulation de la complexité.

Dans XORSAT, les choses se simplifient considérablement. Tout d'abord, nous avons déjà argumenté que le spectre des distances se déduit du spectre des poids du système homogène associé. D'autre part, l'invariance de groupe implique que tous les amas ont le même diamètre. Ce diamètre commun est donné par le plus grand poids représenté dans l'amas contenant $(0, \dots, 0)$. On est ainsi ramené à un problème d'optimisation simple portant sur une solution σ du système homogène. Dans la limite

$\beta \rightarrow -\infty$, les équations de cavité de ce problème d'optimisation s'expriment comme :

$$p_{i \rightarrow a} = \delta_0 \quad \text{si } m_{i \rightarrow a} \neq *, \quad h_{i \rightarrow a} = \sum_{b \in \partial i \setminus a} u_{b \rightarrow i} + 1 \quad \text{si } m_{i \rightarrow a} = *,$$

$$q_{a \rightarrow i} = \delta_0 \quad \text{si } n_{a \rightarrow i} \neq *, \quad u_{a \rightarrow i} = -\mathcal{S} \left[\prod_{\substack{j \in \partial a \setminus i \\ m_{j \rightarrow a} = *}} (-h_{j \rightarrow a}) \right] \min_{\substack{j \in \partial a \setminus i \\ m_{j \rightarrow a} = *}} |h_{j \rightarrow a}| \quad \text{si } n_{a \rightarrow i} = *,$$

avec les notations

$$\beta h_{i \rightarrow a} = \frac{1}{2} \log \frac{p_{i \rightarrow a}(0)}{p_{i \rightarrow a}(1)} \quad \text{et} \quad \beta u_{i \rightarrow a} = \frac{1}{2} \log \frac{q_{a \rightarrow i}(0)}{q_{a \rightarrow i}(1)}. \quad (6.67)$$

Notez que les conditions $m_{i \rightarrow a} \neq *$ et $n_{a \rightarrow i} \neq *$ ne dépendent pas de l'amas considéré.

Une formule similaire donne le diamètre :

$$w_d = \sum_{i, m_i = *}^N \frac{1 + \mathcal{S}(\sum_{a \in i} u_{a \rightarrow i} + 1)}{2} \quad (6.68)$$

Ces équations peuvent aisément être résolues sur un graphe donné, ou dans leur version moyennée sur un arbre infini. La figure 6.6 donne le diamètre réduit $x_1 = w_d/N$ des amas en fonction de la densité de tests $\alpha = M/N$, dans la limite thermodynamique.

À ce stade il paraît nécessaire de clarifier certains concepts. Dans la limite de température nulle, les convictions tendent à devenir déterministes, même quand $m_{i \rightarrow a} = *$. Les champs $h_{i \rightarrow a}$ et les biais $u_{a \rightarrow i}$ deviennent entiers et quantifient le degré d'exigence associé à une conviction. Cependant, bien que de telles convictions ressemblent fort à des avertissements, la tyrannie qu'elles exercent n'est pas aussi forte que celle des véritables avertissements. En effet, les contradictions, manifestées par exemple par des biais incidents $u_{a \rightarrow i}$ de signes opposés, se résolvent par l'optimisation locale du surcoût énergétique, tandis que les contradictions entre avertissements sont irréconciliables. C'est toute la différence entre optimisation et satisfaction de contraintes, exprimée ici à un niveau local. Dans le problème présent, nous avons affaire à un mélange des deux, car la maximisation de la distance est restreinte à l'espace des solutions.

6.3.2 Distances entre amas

En présence d'une multiplicité d'état, l'analyse des propriétés géométriques de l'espace des solutions doit reposer sur une généralisation de l'outil d'énumération des

distances utilisé au chapitre précédent. La fonction « classique » d'énumération des distances :

$$Z(\beta) = \sum_{\sigma, \sigma' \in \mathcal{F}} 2^{-\beta \|\sigma - \sigma'\|}. \quad (6.69)$$

présente l'avantage de se prêter à une analyse ne reposant sur aucune connaissance *a priori* du phénomène de fragmentation, comme en témoigne le succès de la stratégie adoptée au paragraphe 5.2. En revanche, en présence d'une phase fragmentée, il est utile d'introduire la fonction d'énumération généralisée :

$$\mathcal{Z}(\beta, m) = \sum_{c, c'} 2^{-\beta m F(c, c', \beta)} \doteq \sum_{c, c'} \left(\sum_{(\sigma, \sigma') \in c \times c'} 2^{-\beta \|\sigma - \sigma'\|} \right)^m. \quad (6.70)$$

La température inverse interne m permet d'établir une hiérarchie entre les fluctuations internes aux paires d'amas, décrites par les fonctions $F(c, c', \beta)$, et les fluctuations d'amas à amas.

Les limites de température nulle ($\beta \rightarrow \pm\infty$) sont d'un intérêt particulier, car elles permettent d'étudier la statistique des distances minimales et maximales entre amas. Quand on maintient le paramètre $\gamma = \beta m$ fini à mesure que $|\beta|$ tend vers l'infini, on obtient deux fonctions de partition :

$$\mathcal{Z}_{\min}(\gamma) \doteq 2^{\psi_{\min}(\gamma)} = \sum_{c, c'} 2^{-\gamma w_{\min}(c, c')} \quad \text{et} \quad \mathcal{Z}_{\max}(\gamma) \doteq 2^{\psi_{\max}(\gamma)} = \sum_{c, c'} 2^{-\gamma w_{\max}(c, c')}, \quad (6.71)$$

pour $\beta \rightarrow +\infty$ et $-\infty$ respectivement, qui définissent deux fonctions d'énumération des distances *entre amas* : l'une pour les distances minimales

$$w_{\min}(c, c') \doteq \lim_{\beta \rightarrow \infty} F(c, c', \beta) = \min_{(\sigma, \sigma') \in c \times c'} \|\sigma - \sigma'\| \quad (6.72)$$

et l'autre pour les distances maximales

$$w_{\max}(c, c') \doteq \lim_{\beta \rightarrow -\infty} F(c, c', \beta) = \max_{(\sigma, \sigma') \in c \times c'} \|\sigma - \sigma'\|. \quad (6.73)$$

Le choix de température infinie ($\beta = 0$) permet de définir une troisième fonction de moindre importance, qui énumère les paires d'amas en fonction de la distance *typique* entre deux solutions de ces amas.

Le spectre des distances est formé de deux intervalles connexes : d'un côté, un intervalle s'étirant de 0 au diamètre maximal d'un amas ; de l'autre, un intervalle délimité par les distances minimales x_2 et maximales x_3 entre amas. Pour les mêmes raisons que dans le cas des diamètres, ces deux dernières quantités ne s'obtiennent pas

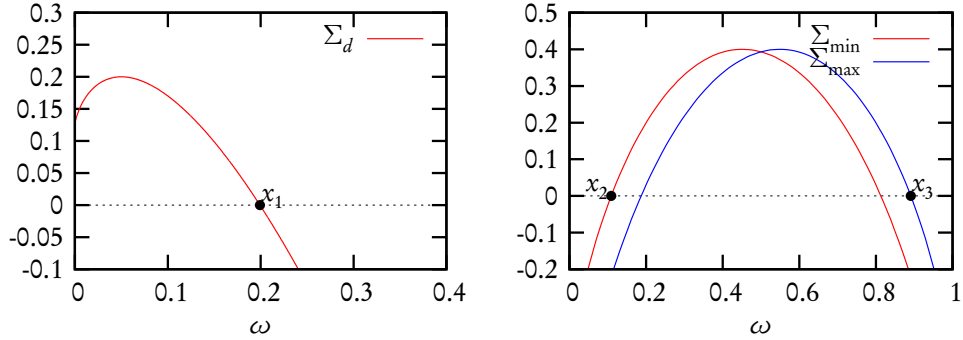


Fig. 6.7: À gauche : fonction d'énumération des diamètres dans le modèle à amas aléatoires ($\alpha = 0,8$ et $p = 0,95$). À droite : fonctions d'énumérations des distances minimales et maximales dans le même modèle. Dans les deux cas, le diamètre maximal, ou les distances minimales et maximales entre amas, s'obtiennent par l'annulation des complexités.

par la simple limite $y \rightarrow \pm\infty$, mais par l'annulation des deux fonctions de complexité (voir la figure 6.7 pour s'en convaincre dans le cas du modèle à amas aléatoires) :

$$\Sigma_{\min}(\omega) = \frac{1}{N} \log \sum_{c,c'} \mathbb{I}(w_{\min}(c, c') = N\omega) \quad \text{et} \quad \Sigma_{\max}(\omega) = \frac{1}{N} \log \sum_{c,c'} \mathbb{I}(w_{\max}(c, c') = N\omega). \quad (6.74)$$

qui se déduisent des potentiels $\psi_{\min}(y)$ et $\psi_{\max}(y)$ par des transformations de Legendre. En pratique, ces potentiels s'écrivent de la même façon que (6.64), la condition de gel simultané (6.63) en moins.

Dans XORSAT, on peut se contenter de considérer les distances minimales et maximales à $(0, \dots, 0)$, que l'on note $w_{\min}(c)$ et $w_{\max}(c)$. Dans les définitions précédentes, les sommes sur les paires d'amas (c, c') peuvent donc être remplacées par de simples sommes sur c . Comme l'ensemble des avertissements non contraignants ne dépend pas de l'amas considéré, on peut fixer une fois pour toute la forme des convictions suivant leur nature. Ainsi, si $m_{a \rightarrow i} = *$, il convient de travailler avec une probabilité de champ $P_{i \rightarrow a}(h_{i \rightarrow a})$, où $h_{i \rightarrow a}$ est défini éq. (6.67). Dans le cas contraire, les probabilités $\pi_{i \rightarrow a}^0 + \pi_{i \rightarrow a}^1 = 1$ suffisent. Nous nous étendons pas sur le détail des équations, qui découlent du formalisme général exposé au §6.1.2. Ces équations sont reproduites dans l'article [MM06b], page 16.

La figure 6.8 fournit un exemple de fonctions de complexité calculées par la méthode de la cavité dans un problème k -XORSAT aléatoire donné.

Quand le problème n'est pas effeuillable, c'est-à-dire quand il n'existe pas de variables de degré 0 ou 1, les amas sont réduits à des singletons, et l'ensemble des variables est gelé. Dans ce cas, les fonctions d'énumération $\psi_{\min}$ et $\psi_{\max}$ se réduisent toutes les deux à la fonction d'énumération classique. Du point de vue des équations

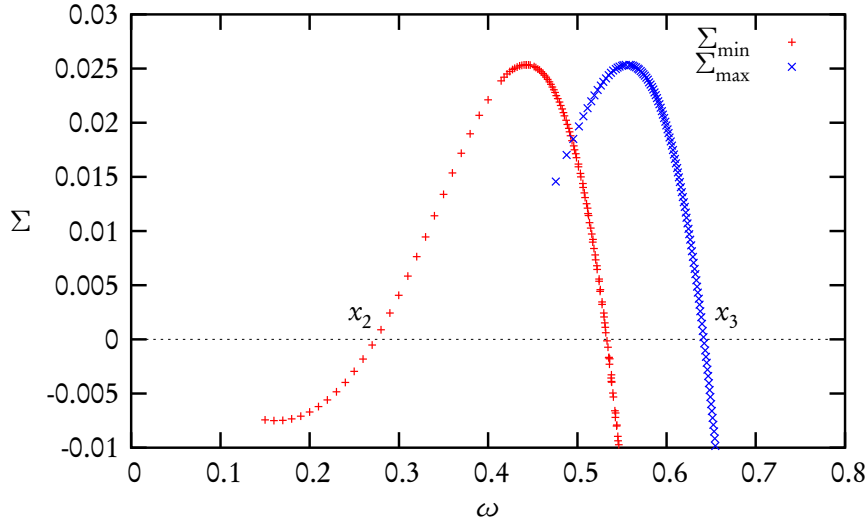


Fig. 6.8: Fonctions d'énumération des poids minimaux et maximaux dans un problème 3-XORSAT aléatoire de taille $N = 10\,000$ et $M = 8\,600$. L'annulation de ces fonctions donne les distances minimales x_2 et maximales x_3 entre amas distincts.

de cavité, la propagation des sondages est alors entièrement décrite par les probabilités d'avertissement $\pi_{i \rightarrow a}^0$, $\pi_{i \rightarrow a}^1$ et $\eta_{a \rightarrow i}^0$, $\eta_{a \rightarrow i}^1$, qui vérifient :

$$\pi_{i \rightarrow a}^0 \propto \prod_{b \in \partial i \setminus a} \eta_{b \rightarrow i}^0, \quad \pi_{i \rightarrow a}^1 \propto 2^{-y} \prod_{b \in \partial i \setminus a} \eta_{b \rightarrow i}^1, \quad (6.75)$$

$$\eta_{a \rightarrow i}^0 = \frac{1 + \prod_{j \in \partial a \setminus i} (2\pi_{j \rightarrow a}^0 - 1)}{2}. \quad (6.76)$$

Avec les identifications $\pi_{i \rightarrow a}^\sigma = p_{i \rightarrow a}(\sigma)$, $\eta_{a \rightarrow i}^\sigma = q_{a \rightarrow i}(\sigma)$, et $y = \beta$, on retrouve exactement les équations à un état unique (4.56), (4.57) intervenant dans le calcul de la fonction d'énumération des poids des mots de codes ($A = 2\mathbb{N}$).

Cette réduction illustre avec clarté une remarque déjà formulée au §5.2.1 : en l'absence de fluctuations au sein des amas, l'hypothèse d'un état unique peut rendre compte de la statistique de ces amas.

x -satisfaisabilité

Les calculs de fonctions d'énumération permettent de remonter au seuil de x -satisfaisabilité. En effet, à α fixé, un problème est x -satisfaisable presque sûrement si et seulement si $x \in [0, x_1] \cup [x_2, x_3]$. La figure 6.9 représente le seuil $\alpha_s(x)$ construit à partir du calcul de ces distances dans le modèle k -XORSAT. Comme dans le modèle à amas aléatoires (figure 5.5), on y observe que le seuil de fossé α_{fos} est supérieur au seuil de fragmentation α_d .

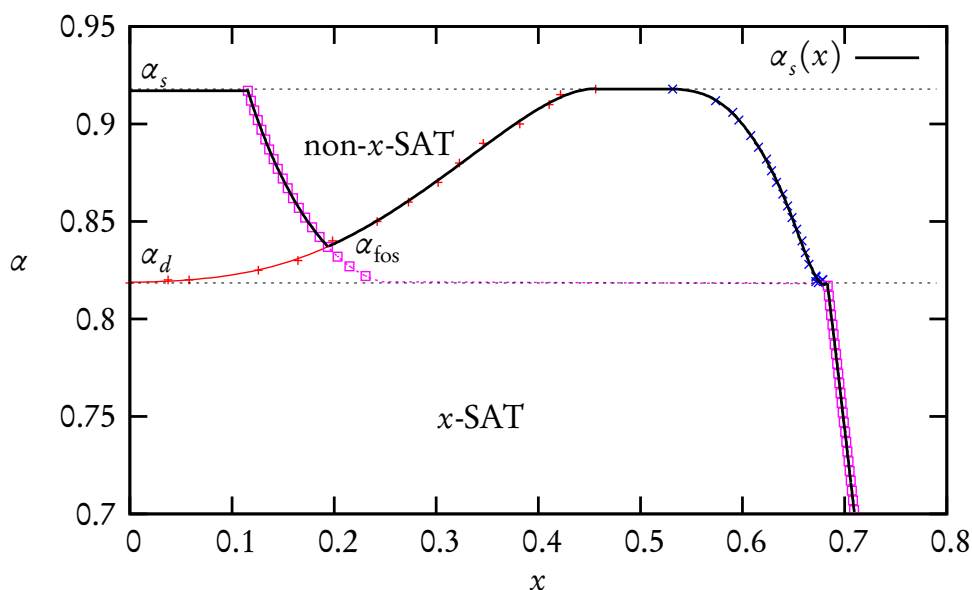


Fig. 6.9: Le seuil de x -satisfaisabilité dans 3-XORSAT aléatoire. Le diamètre d'amas x_1 ($\square$), ainsi que les distances minimales x_2 ($+$) et maximales x_3 ($\times$) entre amas, sont utilisées pour construire le seuil. Comme dans le cas du modèle à amas aléatoires, le fossé apparaît pour une valeur de α supérieure à α_d .

Références

L'extension de la méthode de la cavité à plusieurs états est due à Mézard et à Parisi. À la différence de l'exposé original [MP01], cette méthode est ici formulée comme découlant de la mécanique statistique des points fixes de la propagation des convictions.

Le calcul du seuil de k -SAT est effectué dans [MPZ02, MMZ06], et celui de XORSAT dans [CDMM03, MRTZ03]. Notre définition des modèles étendus généralise la définition proposée par [BZ04, MMW05] dans le cas particulier de k -SAT.

Le calcul des distances, envisagé pour un problème général de satisfaction de contraintes, est mis en pratique sur le problème k -XORSAT, et reprend les résultats de [MM06b].

Conclusion

Distance et ergodicité

Les notions de spectre de distances et de séparabilité sont des outils puissants qui fournissent une information univoque sur la structure géométrique de l'espace des solutions. Néanmoins, elles ne suffisent pas à rendre compte de tous les aspects de cette structure. Si le lien entre les états « purs » identifiés aux solutions de cavité et les composantes connexes peut aisément être explicité dans le problème k -XORSAT grâce à la notion de cœur, il reste problématique dans le cas général. En présence de variables gelées, on peut espérer exploiter celles-ci comme signature des amas. Cependant, on sait que certains amas n'admettent aucune variable gelée : ceux-ci restent appréhendables par les méthodes statistiques, mais ils sont difficilement débuscables individuellement. Et quand bien même on disposerait de toutes les solutions de cavité d'un problème donné, l'étude de leurs corrélations et de leurs recouvrements extrêmes resterait une tâche difficile. Il est cependant envisageable d'analyser la structure des amas à l'aide de méthodes purement locales : à partir de solutions données, on opère des changements microscopiques entraînant, sous la pression des contraintes, des réarrangements de plus ou moins grande ampleur [MS05, MS06b, Zho05, Sem07]. Il paraît naturel que l'altération d'une variable gelée entraîne un changement macroscopique faisant basculer le système d'un amas à l'autre. Quand en revanche la variable n'est pas gelée, la taille du réarrangement quantifie à quel point l'amas est filandreux et rétif au changement : un amas épais implique des réarrangements de faible ampleur, et une bonne résistance à l'altération. Ces notions rappellent les critères de stabilité évoqués au §4.1.6. De tels outils pourraient offrir une alternative algorithmique aux considérations statistiques relatives à l' x -satisfaisabilité.

Notre étude s'est presque entièrement concentrée sur l'espace strict des solutions. Pourtant, les notions d'ergodicité, ainsi que celles d'états purs, devraient être envisagées dans un contexte plus général, où le nombre de contraintes violées serait relaxé et envisagé sur tout son spectre. L'extension de l'étude des distances à des seuils d'énergie arbitraires permettraient de quantifier la hauteur des barrières et d'obtenir ainsi une information plus fine sur la structure en vallées. De telles méthodes « mixtes », où distance et nombre de contraintes violées sont traitées selon le même principe, sont

envisageables dans le contexte des méthodes combinatoires comme dans celui de la cavité.

Erreur dans les codes

L'étude en grandes déviations des propriétés de décodage sur le canal d'effacement nous a permis de mettre en évidence un phénomène de transition de phase pour les événements rares. Alors que la physique statistique d'équilibre s'intéresse exclusivement aux propriétés typiques des systèmes, cet exemple nous montre que les grandes déviations peuvent laisser apparaître un diagramme de phase plus riche, invisible au simple niveau typique. Dans le cas des codes de correction d'erreur, l'approche atypique est justifiée par le besoin de contrôler la fréquence des événements rares. Tandis qu'en physique le nombre de degrés de liberté est généralement commensurable au nombre d'Avogadro $N_A = 6 \cdot 10^{23}$, rendant dérisoire le rôle des grandes déviations, la taille des chaînes utilisées pour le codage dépasse rarement $N = 10^5$. Toutefois, il peut arriver que l'étude des événements rares soit justifiée dans l'étude des phénomènes naturels [Ell85, Ell95] : par exemple, l'adaptation d'un réseau de rigidité aux contraintes mécaniques [BBS05, RB06] rentre dans un tel cadre. Parallèlement, l'étude des grandes déviations dans le contexte de la physique statistique hors-équilibre a engendré une activité importante depuis une quinzaine d'années [ES02, BBDR05].

Deux transitions de phases interviennent dans le calcul de la probabilité d'erreur des codes LDPC sous le canal d'effacement : à bruit élevé, l'erreur est majoritairement dominée par le nombre de bits effacés, tandis qu'à bruit plus faible, la direction du bruit devient importante. Lors d'une seconde transition, cette direction de bruit se polarise vers le mot de code le plus proche. Ainsi, l'environnement géométrique joue un rôle déterminant dans la cause de l'erreur, mais la seule connaissance de la distance minimale ne devient prépondérante que dans la phase de très faible bruit. Cette observation suggère la nécessité d'une étude plus poussée de la structure géométrique, où la notion de direction serait prise en compte.

Alors que l'analyse du décodage optimal amène à considérer le spectre des mots de codes, le succès du décodage itératif est, dans le cas du canal d'effacement, corrélatif de la présence de sous-parties d'arrêts. L'extension de ce critère à un canal plus général conduit à considérer les *pseudo mots de code*, qui généralisent les sous-parties d'arrêt. L'étude statistique de ces objets et leur incidence sur les performances de décodage constituent un défi de première importance en théorie de l'information. Notons au passage que les questions relatives au lien entre distances minimales et performances algorithmiques ne sont pas l'apanage des codes LDPC, et sont également importantes dans le contexte des Turbo-codes [BGT93], qui constituent actuellement la meilleure alternative aux codes LDPC en termes de fiabilité. Enfin, mentionnons une approche alternative, également inspirée de la physique, qui explique l'erreur dans les codes par

des « instantons » [SCCV05], et fait le lien avec les pseudo mots de code.

Dans l'optique de l'amélioration des codes existants, une direction difficile, mais qui pourrait s'avérer prometteuse, consiste à combler le fossé qui sépare la performance itérative de la performance optimale. Bien qu'il existe des constructions de codes dont le seuil itératif approche de très près la borne de Shannon, ces codes présentent dans le même temps de très médiocres propriétés d'erreur à faible bruit. Inversement, les codes à grand degré moyen ont de bonne propriété d'erreur mais admettent un très mauvais seuil itératif de décodabilité, bien que leur seuil optimal sature la borne de Shannon. L'élaboration de codes praticables et universellement performants, près de la borne de Shannon comme dans la limite de bruit nul, constituerait un progrès considérable. Selon une interprétation physique, l'échec du décodage itératif est causé par la présence d'une phase « vitreuse », qui bloquerait l'algorithme dans des états métastables d'énergie sous-optimale⁸. Il est donc envisageable d'exploiter les techniques inspirées de la propagation des sondages afin d'aborder ce problème. En pratique cependant la tâche s'avère difficile, en raison des propriétés de symétrie des tests de parité. Un point de vue complémentaire à l'interprétation vitreuse, proposé dans le cas du canal d'effacement [MMU05], relie les décodages optimal et itératif par le truchement d'une transformation de Maxwell, en situant le mot de code optimal dans une poche non-convexe de l'ensemble micro-canonique. En suivant une approche similaire, l'optimalité de BP a pu être prouvée [MT06] pour une classe de codes servant à la transmission de messages provenant de canaux multiples (les codes CDMA : *code-division multiple-access*). Ces codes diffèrent des codes LDPC car ils sont entièrement connectés (les facteurs font intervenir toutes les variables), et non-linéaires.

La méthode de la cavité

La méthode de la cavité a occupé une place importante dans ce travail de thèse. Tout d'abord, la vérification de sa validité en conjonction avec l'hypothèse de la multiplicité d'états a été l'une des motivations principales conduisant à l'étude de la fragmentation et de l' x -satisfaisabilité. Ensuite, dans un mouvement de retour, nous l'avons utilisée afin de dériver le seuil de x -satisfaisabilité dans XORSAT. Enfin, elle est à la base des méthodes de grandes déviations employées dans les calculs d'erreur dans les codes.

L'exposé a présenté la méthode de la cavité comme un Ansatz exact sur les arbres, et comme une approximation asymptotiquement correcte sur les graphes aléatoires dilués, sous certaines conditions. Le domaine d'application de la méthode est pourtant potentiellement beaucoup plus large. Par exemple, on connaît son efficacité dans les

⁸La distance joue ici le rôle d'énergie.

modèles de verre de spin en champ moyen, tels le modèle de Sherrington-Kirkpatrick (SK) [SK75] ou le modèle p -spin [GM84], qui sont entièrement connectés, mais où chaque interaction est de faible amplitude. Il est fort probable qu'une large classe de problèmes intermédiaires, entre les graphes dilués et les graphes entièrement connectés, puisse être appréhendables par la méthode de la cavité.

Le succès de la méthode de la cavité peut se mesurer à sa fécondité en théorèmes mathématiques. Depuis la solution de Parisi au modèle SK, les méthodes des répliques et de la cavité ont activement été étudiées par les mathématiciens [Tal03]. Talagrand, en se basant sur les travaux de Guerra, a ainsi pu rigoureusement établir la validité des schémas de brisure de symétrie des répliques pour le modèle SK [Tal06] et pour le modèle p -spins [Tal00]. Mais le succès mathématique de la méthode de la cavité ne se limite pas aux modèles entièrement connectés. Ainsi, les prédictions des physiciens sur le modèle d'appariement aléatoire [MP86], originellement formulées dans le cadre des répliques, ont pu être confirmées par Aldous [Ald01] à l'aide de la méthode de la cavité. L'utilisation de la méthode d'interpolation de Guerra a également permis de dériver une borne supérieure exacte sur l'estimation de l'entropie dans une large classe de problèmes dilués, dont k -SAT [FL03], où cette borne est précisément la prédiction de la cavité, supposée exacte. Dans le même registre, l'unicité de la solution de cavité à état unique, ainsi que sa validité, ont été prouvées pour les petites densités de contraintes en utilisant un critère de reconstructibilité (ce résultat est l'objet d'un article à paraître par Dembo et Montanari). Du côté de la théorie du codage, l'analyse par l'«évolution des densités» de messages dans le décodage itératif [RU01] a permis de dériver le seuil dynamique des codes LDPC, indépendamment des méthodes «physiques». Pour ce qui est du seuil de décodabilité optimale, il a été établi rigoureusement dans le cas du canal d'effacement [MMRU04] et il existe, pour le canal général, une borne supérieure [Mon05] dérivée à l'aide des polynômes d'interpolation de Guerra. On le voit, la transformation des prédictions de la cavité en théorèmes est la source d'une importante activité mathématique, et on peut estimer que cette direction de recherche est appelée à perdurer dans l'avenir.

D'un point de vue plus pratique, la propagation des convictions a fait preuve de son efficacité sur les codes LDPC. Son extension aux états multiples, incarnée par la propagation des sondages, a permis un progrès important dans la résolution des problèmes de satisfaction de contraintes aléatoire difficiles [BMWZ02]. Cependant, dans les deux cas, les algorithmes naïfs de propagation de messages se heurtent à des difficultés dès que le graphe n'est plus aléatoire. En cause, les petites boucles, absentes des graphes aléatoires mais fréquentes dans les problèmes réels, requièrent un traitement spécial au delà de l'approximation des arbres. La question de savoir comment mettre en œuvre ces modifications reste largement ouverte, et mériterait d'être étudiée plus avant.

Perspectives

On assiste depuis quelques années à une convergence des méthodes dans les domaines de l'inférence, de la complexité algorithmique, de la théorie de la communication et de la physique statistique. Au cœur de ces méthodes, les techniques de passage de message sont basées sur l'idée qu'une tâche globale d'inférence ou d'optimisation peut être distributivement résolue par une somme de tâches locales n'interférant qu'à courte portée. Bien que ce principe soit mis en défaut par le phénomène de fragmentation dans certains problèmes complexes, la propagation des sondages offre une voie de sortie, en compilant l'information des états purs, qui sont autant de solutions au problème local, sous la forme de messages généralisés. Outre les domaines sus-mentionnés, ces méthodes ont trouvé des champs d'application dans l'analyse de données, notamment sur le problème de regroupement des données (*data clustering*) [FD07], et en inférence dans des contextes biologiques : par exemple l'apprentissage du perceptron [BZ06] ou l'inférence sur les réseaux de régulation de gènes [CLP⁺06a, CLP⁺06b, MPWZ07].

La fragmentation sert-elle à quelque chose ? Dans la résolution des problèmes de satisfaction de contraintes, la fragmentation représente évidemment un obstacle. En revanche, du point de vue de la théorie de l'information, elle joue un rôle déterminant, car elle met en œuvre un principe de discrimination. Nous l'avons vu en détail dans les codes de correction d'erreur, où les amas ont la particularité d'être ponctuels. Mais cette remarque reste pertinente même quand les amas sont épais : il a par exemple été montré [BBCZ05] que les amas de k -SAT peuvent efficacement représenter les messages d'une source dans un schéma de compression avec perte.

L'analyse récente [SBSB06, TSBB06] d'expériences sur l'activité jointe de neurones ganglionnaires de la rétine de vertébrés suggère que ce même principe est à l'œuvre dans les premières couches du traitement neuronal de l'information visuelle. Selon l'interprétation de cette expérience, les cellules ganglionnaires agissent comme un filtre, en réduisant l'ensemble des stimuli possibles à un nombre réduit d'états, ou vallées, prescrites par un modèle de verre de spins sous-jacent. Ce filtrage opère précisément un principe de compression avec perte, en réduisant l'entropie de la source à un nombre réduit de traits caractéristiques. Cette réduction dimensionnelle engendre par ailleurs des messages robustes, car le passage d'une vallée à l'autre implique l'apport d'un bruit important, de manière tout-à-fait analogue avec ce qu'on observe dans les codes de correction d'erreur. À la lumière de ces observations, il semble que la fragmentation dans les modèles graphiques ne doit pas être simplement perçue comme un obstacle, mais aussi comme un outil permettant de créer une structure aux propriétés discriminatoires, à l'aide d'interactions purement locales. Bien que cette idée soit banale en physique statistique, où il est bien connu que l'ordre ferromagnétique est induit par les interactions à courte portée, elle revêt un intérêt nouveau dans un

contexte désordonné, où le nombre de vallées est exponentiel et permet un processus de sélection sans pour autant réduire l'information à presque rien⁹.

Cet essaimage des techniques et concepts issus de la théorie des verres dans de nombreuses branches de la science, et en particulier dans le traitement biologique de l'information, amorce à notre avis une voie potentiellement fructueuse de recherche.

⁹Un modèle ferromagnétique simple tel que le modèle d'Ising ne contient qu'un *bit* d'information, car seuls deux états sont possibles.

Articles

“Clustering of Solutions in the Random
Satisfiability Problem”

Phys. Rev. Lett. **94** 197205 (2005)

Clustering of Solutions in the Random Satisfiability Problem

M. Mézard,¹ T. Mora,¹ and R. Zecchina²

¹*Laboratoire de Physique Théorique et Modèles Statistiques, bâtiment 100, Université Paris-Sud, F-91405 Orsay, France.*

²*Abdus Salam International Center for Theoretical Physics, Strada Costiera 11, 34100 Trieste, Italy*

(Received 18 February 2005; published 19 May 2005)

Using elementary rigorous methods we prove the existence of a clustered phase in the random K -SAT problem, for $K \geq 8$. In this phase the solutions are grouped into clusters which are far away from each other. The results are in agreement with previous predictions of the cavity method and give a rigorous confirmation to one of its main building blocks. It can be generalized to other systems of both physical and computational interest.

DOI: 10.1103/PhysRevLett.94.197205

PACS numbers: 75.10.Nr, 75.40.Mg

Constraint satisfaction problems (CSPs) provide one of the main building blocks for complex systems studied in computer science, information theory, and statistical physics, and may even turn out to be important in the statistical studies of biological networks. Typically, they involve a large number of discrete variables, each one taking a finite number of values, and a set of constraints: each constraint involves a few variables, and forbids some of their joint assignments. A simple example is the q coloring of a graph, where one should assign to each vertex of the graph a color in $\{1, \dots, q\}$, in such a way that two vertices related by an edge have different colors. In the case $q = 2$, this is nothing but the zero temperature limit of an antiferromagnetic problem, which is known to display a spin glass behavior when the graph is frustrated and disordered. CSPs also appear naturally in the studies of structural glasses [1] and rigidity percolation [2].

Given an instance of a CSP, one wants to know whether there exists a solution, that is, an assignment of the variables which satisfies all the constraints (e.g., a proper coloring). When it exists the instance is called satisfiable, and one wants to find a solution. Most of the interesting CSPs are NP-complete problems: in the worst case the number of operations needed to decide whether an instance is SAT or not is expected to grow exponentially with the number of variables. But recent years have seen an upsurge of interest in the theory of typical case complexity, where one tries to identify random ensembles of CSPs which are hard to solve, and the reason for this difficulty. Random ensembles of CSPs are also of great theoretical and practical importance in communication theory: some of the best error correcting codes (the so-called low density parity check codes) are based on such constructions [3,4].

The archetypical example of CSP is the satisfiability problem (SAT). This is a core problem in computational complexity: it is the first one to have been shown to be an NP-complete problem [5], and since then thousands of problems have been shown to be computationally equivalent to it. Yet it is not so easy to find difficult instances. The main ensemble which has been used for this goal is the

random K -satisfiability (K -SAT) ensemble. The variables are N binary variables—Ising spins— $\vec{\sigma} = \{\sigma_i\} \in \{-1, 1\}^N$. The constraints are called K -clauses. Each of them involves K distinct spin variables, randomly chosen with uniform distribution, and it forbids one configuration of these spins, randomly chosen among the 2^K possible ones. A set of M clauses defines the problem. This corresponds to generating a random logical formula in conjunctive normal form, which is a very generic problem appearing in logic. K -SAT can also be written as the problem of minimizing a spin glasslike energy function which counts the number of violated clauses and in this respect random K -SAT is seen as a prototypical diluted spin glass [6]. Here we shall keep to the most interesting case $K \geq 3$ (for $K = 2$ the problem is polynomial).

In the recent years random K -SAT has attracted much interest in computer science and in statistical physics [7–10]. The interesting limit is the thermodynamic limit $N \rightarrow \infty$, $M \rightarrow \infty$ at fixed clause density $\alpha = M/N$. Its most striking feature is certainly its sharp threshold. It is strongly believed that there exists a phase transition for this problem: Numerical and heuristic analytical arguments are in support of the so-called *satisfiability threshold conjecture*: *There exists $\alpha_c(K)$ such that, with high probability, if $\alpha < \alpha_c(K)$, a random instance is satisfiable; if $\alpha > \alpha_c(K)$, a random instance is unsatisfiable.* Throughout this Letter “with high probability” (w.h.p.) means with a probability going to one in the $N \rightarrow \infty$ limit. Although this conjecture remains unproven, Friedgut has come close to it by establishing the existence of a nonuniform sharp threshold [11]. A lot of effort has been devoted to understanding this phase transition. This is interesting not only from the physics point of view but also from the computer science one, because the random instances with α close to α_c are the hardest to solve. The most important rigorous results so far are bounds for the threshold $\alpha_c(K)$. The best upper bounds were derived using first moment methods [12,13]. Lower bounds can be found by analyzing some algorithms which find SAT assignments [14,15], but recently a new method, based on second moment methods, has found better and

algorithm-independent lower bounds [16,17]. Using these bounds, it was shown that $\alpha_c(K)$ scales as $2^K \ln(2)$ when $K \rightarrow \infty$.

On the other hand, some claim that the cavity method, which is a powerful tool from the statistical physics of disordered systems [18], can be used to compute the exact value of the threshold [19–21], giving for instance $\alpha_c(3) \approx 4.2667$. It is a nonrigorous method but the self-consistency of its results have been checked by a “stability analysis” [21–23], and it also led to the development of a new algorithmic strategy, “survey propagation,” which can solve very large instances at clause densities which are very close to the threshold (e.g., $N = 10^6$ and $\alpha = 4.25$).

The main hypothesis on which the cavity analysis of random K -satisfiability relies is the existence, in a region of clause density $[\alpha_d, \alpha_c]$ close to the threshold, of an intermediate phase called the “hard-SAT” phase. In this phase the set S of solutions (a subset of the vertices in the N -dimensional hypercube) is supposed to split into many disconnected clusters $S = S_1 \cup S_2 \cup \dots$. If one considers two solutions X, Y in the same cluster S_j , it is possible to walk from X to Y (staying in S) by flipping at each step a finite number of spins. If on the other hand X and Y are in different clusters, in order to walk from X to Y (staying in S), at least one step will involve an extensive number (i.e., $\propto N$) of spin flips. This clustered phase is held responsible for trapping many local search algorithms into nonoptimal metastable states [24]. This phenomenon is not exclusive to random K -SAT. It is also predicted to appear in many other hard satisfiability and optimization problems such as coloring [25,26] or the multi-index matching problem [27], and corresponds to a “one step replica symmetry breaking” phase in the language of statistical physics. It is also a crucial limiting feature for decoding algorithms in some error correcting codes [28]. So far, the only CSP for which the existence of the clustering phase has been established rigorously is the simple polynomial problem of random exclusive-OR-SAT (XOR-SAT) [29,30]. In other cases it is an hypothesis, the self-consistency of which is checked by the cavity method.

In this Letter we provide rigorous arguments which show the existence of the clustering phenomenon in random K -SAT, for large enough K , in some region of α included in the interval $[\alpha_d(K), \alpha_c(K)]$ predicted by the statistical physics analysis. Our result is not able to confirm all the details of this analysis but it provides strong evidence in favor of its validity.

Given an instance F of random K -satisfiability, we define a SAT- x -pair as a pair of assignments $(\vec{\sigma}, \vec{\tau}) \in \{-1, 1\}^{2N}$, which both satisfy F , and which are at a Hamming distance $d_{\sigma\tau} \equiv \sum_{i=1}^N (1 - \sigma_i \tau_i)/2$ specified by x as follows:

$$d_{\sigma\tau} \in [Nx - \epsilon(N), Nx + \epsilon(N)] \quad (1)$$

Here x is the normalized distance between the two con-

figurations, which we keep fixed as N and d go to infinity. The resolution $\epsilon(N)$ must be such that $\lim_{N \rightarrow \infty} \epsilon(N)/N = 0$, but its precise form is unimportant for our large N analysis. One can choose for instance $\epsilon(N) = \sqrt{N}$.

We call x -satisfiable a formula for which such a pair of solutions exists. Our study mimics the usual steps which are taken in rigorous studies of K -SAT, but taking pairs of assignments at a fixed distance instead of single assignments.

We first formulate the x -satisfiability threshold conjecture: For all $K \geq 2$ and for all $x, 0 < x < 1$, there exists an $\alpha_c(K, x)$ such that, w.h.p., if $\alpha < \alpha_c(K, x)$, a random K -CNF is x -satisfiable; if $\alpha > \alpha_c(K, x)$, a random K -CNF is x -unsatisfiable, which generalizes the usual satisfiability threshold conjecture (obtained for $x = 0$). We shall find explicitly below two functions, $\alpha_{LB}(K, x)$ and $\alpha_{UB}(K, x)$, which give lower and upper bounds for α for x -satisfiability at a given value of K . Numerical computations of these bounds show that $\alpha(K, x)$ is nonmonotonous as a function of x for $K \geq 8$, as illustrated in Fig. 1. This in turn shows that, for K large enough and in some well chosen interval of α below the satisfiability threshold, SAT- x -pairs exist for x close to 0 ($\vec{\sigma}$ and $\vec{\tau}$ in the same cluster) and x close to 0.5 ($\vec{\sigma}$ and $\vec{\tau}$ in different clusters),

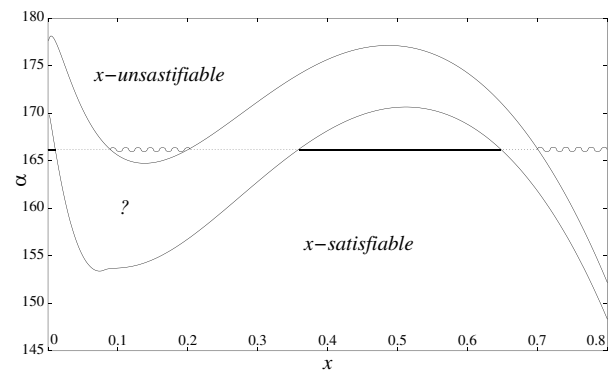


FIG. 1. Lower and upper Bounds for the x -satisfiability threshold $\alpha_c(K = 8, x)$. The upper curve is obtained by the first moment method. Above this curve there exists no SAT- x -pair, w.h.p. The lower curve is obtained by the second moment method. Below this curve there exists a SAT- x -pair, w.h.p. For values of α lying between 164.735 and 170.657, these bounds guarantee the existence of a clustering phenomenon. The horizontal line gives an example of this phenomenon for $\alpha = 166.1$. We exhibit the successive phases as one varies x : x -satisfiable regions are represented by a thick solid line, x -unsatisfiable regions by a wavy line, and “donot know” regions by a dotted line. The x -satisfiable region near $x = 0$ corresponds to intra-cluster pairs, whereas the x -satisfiable region around $x = 0.5$ corresponds to intercluster pairs. In this example, the intermediate x -unsatisfiable region around $x \sim 0.13$ shows the existence of a “gap” between clusters. We recall that the best refined lower and upper bounds for the satisfiability threshold $\alpha_c(K = 8)$ from [13,17] are, respectively, 173.253 and 176.596. The cavity prediction is 176.543 [21].

but there is an intermediate x region where they do not exist. Figure 1 shows an explicit example of this scenario for a particular value of α .

In what follows we first establish a rigorous and explicit upper bound using a simple first moment method. Subsequently, we provide a (numerical) lower bound using a second moment method [16,17]. Both results are based on elementary probabilistic techniques which could be generalized to other physical systems or random combinatorial problems.

Upper bound: the first moment method.—We use the fact that, when Z is a non-negative random variable,

$$\mathbf{P}(Z \geq 1) \leq \mathbf{E}(Z). \quad (2)$$

Given a formula F , we take $Z(F)$ to be the number of pairs of solutions at fixed distance [with resolution $\epsilon(N)$]:

$$Z(F) = \sum_{\vec{\sigma}, \vec{\tau}} \delta\left(\frac{d_{\sigma\tau}}{N} \approx x\right) \delta(\vec{\sigma}, \vec{\tau} \in S(F)), \quad (3)$$

where $S(F)$ is the set of solutions to F . Throughout this Letter $\delta(A)$ is an indicator function, equal to one if the statement A is true, and to 0 otherwise. Since $Z(F) \geq 1$ is equivalent to “ F is x -satisfiable,” (3) gives an upper bound for the probability of x -satisfiability. The expected value of the double sum over the choice of a random F is

$$\mathbf{E}(Z(F)) = 2^N \binom{N}{Nx} \mathbf{E}[\delta(\vec{\sigma}, \vec{\tau} \in S(c))]^M. \quad (4)$$

We have used $\delta(\vec{\sigma}, \vec{\tau} \in S(F)) = \prod_c \delta(\vec{\sigma}, \vec{\tau} \in S(c))$, where c denotes the clauses, and the fact that clauses are drawn independently. The expectation $\mathbf{E}[\delta(\vec{\sigma}, \vec{\tau} \in S(c))]$ is equal to $1 - 2^{1-K} + 2^{-K}(1-x)^K$ (there are only two realizations of the clause among 2^K that do not satisfy c unless the two configurations overlap exactly on the domain of c).

In the thermodynamic limit, $\ln \mathbf{E}(Z(F))/N \rightarrow \Phi_1(x, \alpha)$, where

$$\Phi_1(x, \alpha) = \ln 2 + H_2(x) + \alpha \ln[1 - 2^{-K}[2 - (1-x)^K]],$$

where $H_2(x) = -x \ln x - (1-x) \ln(1-x)$ is the two-state entropy function. This gives the upper bound

$$\alpha_{\text{UB}}(K, x) = -\frac{\ln 2 + H_2(x)}{\ln[1 - 2^{1-K} + 2^{-K}(1-x)^K]}. \quad (5)$$

Lower bound: the second moment method.—We use the fact that, when Z is a non-negative random variable,

$$\mathbf{P}(Z > 0) \geq \frac{\mathbf{E}(Z)^2}{\mathbf{E}(Z^2)}. \quad (6)$$

However, using this formula with Z equal to the number of solutions fails, and one must instead use a weighted sum [16]. We follow the strategy recently developed in [17], which we generalize to SAT- x -pairs by taking

$$Z(F) = \sum_{\vec{\sigma}, \vec{\tau}} \delta\left(\frac{d_{\sigma\tau}}{N} \approx x\right) \prod_c W(\vec{\sigma}, \vec{\tau}, c). \quad (7)$$

$W(\vec{\sigma}, \vec{\tau}, c)$ is a weight associated with the clause c , given the couple $(\vec{\sigma}, \vec{\tau})$, and is defined as follows: Suppose that c is satisfied by n_σ among the K $\vec{\sigma}$ variables involved in c , and by n_τ among the K $\vec{\tau}$ variables. Call n_0 the number of common values between the $\vec{\sigma}$ and $\vec{\tau}$ variables involved in c . Then define

$$W(\vec{\sigma}, \vec{\tau}, c) = \begin{cases} \lambda^{n_\sigma + n_\tau} \nu^{n_0} & \text{if } n_\sigma > 0 \text{ and } n_\tau > 0, \\ 0 & \text{otherwise.} \end{cases} \quad (8)$$

Note that with this definition of Z the choice $\lambda = 1$, $\nu = 1$ simply yields the number of solutions (3).

Let us now compute the first two moments of Z ([31]):

$$\mathbf{E}(Z) = 2^N \binom{N}{Nx} \left[f_1^{(\lambda, \nu)}(x) \right]^M, \quad (9)$$

where $f_1^{(\lambda, \nu)}(x) = \mathbf{E}(W(\vec{\sigma}, \vec{\tau}, c))$ can be calculated by simple combinatorics (via multinomial sums). To compute $\mathbf{E}(Z^2)$, we sum over four spin configurations $\vec{\sigma}, \vec{\tau}, \vec{\sigma}', \vec{\tau}'$. Symmetry allows us to fix $\sigma_i = 1$. Let $Na(t, s, t')$ be the number of sites i such that $\tau_i = t$, $\sigma'_i = s'$, and $\tau'_i = t'$ (where $t, s, t' \in \{\pm 1\}$). It turns out that the term of the sum depends only on these eight numbers $a(\pm 1, \pm 1, \pm 1)$. We collect them into a vector $\mathbf{a}$ and get

$$\mathbf{E}(Z^2) = 2^N \int_V d\mathbf{a} \frac{N!}{\prod_{t, s', t'} (Na(t, s', t'))!} [f_2^{(\lambda, \nu)}(\mathbf{a})]^M, \quad (10)$$

where $f_2^{(\lambda, \nu)}(\mathbf{a}) = \mathbf{E}(W(\vec{\sigma}, \vec{\tau}, c)W(\vec{\sigma}', \vec{\tau}', c))$ can be calculated by simple combinatorics in the same way as f_1 . The integration set V is a 5-dimensional simplex taking into account the normalization $\sum_{t, s', t'} a(t, s', t') = 1$ and the two constraints: $d_{\sigma\tau}/N \approx x$, $d_{\sigma'\tau'}/N \approx x$.

A saddle point evaluation of Eq. (10) gives, for $N \rightarrow \infty$,

$$\frac{\mathbf{E}(Z)^2}{\mathbf{E}(Z^2)} \geq C_0 \exp\left(-N \max_{\mathbf{a} \in V} \Phi_2(\mathbf{a})\right), \quad (11)$$

where C_0 is a constant depending on K and x , and

$$\begin{aligned} \Phi(\mathbf{a}) = & H_8(\mathbf{a}) - \ln 2 - 2H_2(x) + \alpha \ln f_2^{(\lambda, \nu)}(\mathbf{a}) \\ & - 2\alpha \ln f_1^{(\lambda, \nu)}(x), \end{aligned} \quad (12)$$

with $H_8(\mathbf{a}) = -\sum_{t, s', t'} a(t, s', t') \ln a(t, s', t')$. In general $\max_{\mathbf{a} \in V} \Phi(\mathbf{a})$ is non-negative and one must choose appropriate weights $W(\vec{\sigma}, \vec{\tau}, c)$ in such a way that $\max_{\mathbf{a} \in V} \Phi(\mathbf{a}) = 0$. We notice that at the particular point $\mathbf{a}^*$ where $(\vec{\sigma}, \vec{\tau})$ is uncorrelated with $(\vec{\sigma}', \vec{\tau}')$, we have $\Phi(\mathbf{a}^*) = 0$. We fix the parameters λ and μ defining the weights (8) in such a way that $\mathbf{a}^*$ is a local maximum of Φ . This gives two algebraic equations in λ and ν which have a unique solution $\lambda > 0$, $\nu > 0$. Fixing λ and ν to these

values, α_{LB} is the largest value of α such that the local maximum at $\mathbf{a}^*$ is a *global* maximum, i.e., such that there exists no $\mathbf{a} \in V$ with $\Phi(\mathbf{a}) > 0$:

$$\alpha_{\text{LB}}(K, x) = \inf_{\mathbf{a} \in V} \frac{\ln 2 + 2H_2(x) - H_8(\mathbf{a})}{\ln f_2^{(\lambda, \nu)}(\mathbf{a}) - 2 \ln f_1^{(\lambda, \nu)}(x)}. \quad (13)$$

We devised several numerical strategies to evaluate $\alpha_{\text{LB}}(K, x)$. The implementation of Powell's method starting from each point of a grid of size $\mathcal{N}^5$ ($\mathcal{N} = 10, 15, 20$) on V turned out to be the most efficient and reliable. The results are given by Fig. 1 for $K = 8$, the smallest K such that the clustering conjecture is confirmed. We found a clustering phenomenon for all the values of $K \geq 8$ that we checked, and in fact the relative difference $[\alpha_{\text{UB}}(K, x) - \alpha_{\text{LB}}(K, x)]/\alpha_{\text{LB}}(K, x)$ seems to go to zero at large K .

We have shown a simple probabilistic argument which shows rigorously the existence of a clustered hard-SAT phase. The prediction from the cavity method is in fact a weaker statement. It can be stated in terms of the overlap distribution function $P(x)$, which is the probability, when two SAT assignments are taken randomly (with uniform distribution), that their distance is given by x . The cavity method finds that this distribution has a support concentrated on two values: a small value x_1 , close to zero, gives the characteristic "radius" of a cluster; a larger value x_0 gives the characteristic distance between clusters. This does not imply that there exists no pair of solutions for values of x distinct from x_0, x_1 ; it just means that such pairs are exponentially less numerous than the typical ones. Our rigorous result shows that in fact there exists a true gap in x , with no SAT- x -pairs, at least for $K \geq 8$. More sophisticated moment computations might allow to get some results for smaller values of K . Still the conceptual simplicity of our computation makes it a useful tool for proving similar phenomena in other systems of physical or computational interests, like for instance the graph-coloring (anti-ferromagnetic Potts) problem.

This work has been supported in part by the EC through the network MTR 2002-00319 "STIPCO" and the FP6 IST consortium "EVERGROW."

-
- [1] M. Sellitto, G. Biroli, and C. Toninelli, *Europhys. Lett.* **69**, 496 (2005).
 - [2] J. Barré *et al.*, cond-mat/0408385 [*Phys. Rev. Lett.* (to be published)].
 - [3] Robert G. Gallager, *Information Theory and Reliable Communication* (Wiley, New York, 1968).
 - [4] David J.C. MacKay, *Information Theory, Inference & Learning Algorithms* (Cambridge University Press, Cambridge, 2002).

- [5] Stephen Cook, in *Proceedings of the Third Annual ACM Symposium on Theory of Computing, Shaker Heights, Ohio, United States*, (ACM Press, New York, 1971), p. 151.
- [6] R. Monasson and R. Zecchina, *Phys. Rev. E* **56**, 1357 (1997).
- [7] Special issue on Frontiers in Problem Solving: Phase Transitions and Complexity, edited by T. Hogg, B.A. Huberman, and C. Williams [*Artif. Intell.* **81**, 1 (1996)].
- [8] Special Issue on NP-hardness and Phase transitions, edited by O. Dubois, R. Monasson, B. Selman, and R. Zecchina [*Theor. Comput. Sci.* **265** 1 (2001)].
- [9] S. Kirkpatrick and B. Selman, *Science* **264**, 1297 (1994).
- [10] R. Monasson, R. Zecchina, S. Kirkpatrick, B. Selman, and L. Troyanski, *Nature (London)* **400**, 133 (1999).
- [11] E. Friedgut, *J. Am. Math. Soc.* **12**, 1017 (1999).
- [12] L.M. Kirousis, E. Kranakis, and D. Krizanc, School of Computer Science-Carleton University Technical Report No. TR-96-09, 1996 (unpublished).
- [13] O. Dubois and Y. Boufkhad, *Journal of Algorithms* **24**, 395 (1997).
- [14] M.-T. Chao and J. Franco, *Information Sciences (NY)* **51**, 289 (1990).
- [15] A.M. Frieze and S. Suen, *Journal of Algorithms* **20**, 312 (1996).
- [16] D. Achlioptas and C. Moore, *Proc. Foundations of Computer Science* (2002).
- [17] D. Achlioptas, Y. Peres, *J. Am. Math. Soc.* **17**, 947 (2004).
- [18] M. Mézard and G. Parisi, *J. Stat. Phys.* **111**, 1 (2003).
- [19] M. Mézard and R. Zecchina, *Phys. Rev. E* **66**, 056126 (2002).
- [20] M. Mézard, G. Parisi, and R. Zecchina, *Science* **297**, 812 (2002).
- [21] S. Mertens, M. Mézard, and R. Zecchina, "Threshold values of Random K -SAT from the cavity method" (to be published).
- [22] A. Montanari and F. Ricci-Tersenghi, *Eur. Phys. J. B* **B33**, 339 (2003).
- [23] A. Montanari, G. Parisi, and F. Ricci-Tersenghi, *J. Phys. A* **37**, 2073 (2004).
- [24] G. Semerjian and R. Monasson, *Proceedings of the SAT 2003 Conference*, edited by E. Giunchiglia and A. Tacchella, *Lect. Notes Comput. Sci.* Vol. 2919 (Springer), New York, New York, 2004), p. 120.
- [25] R. Mulet, A. Pagnani, M. Weigt, and R. Zecchina, *Phys. Rev. Lett.* **89**, 268701 (2002).
- [26] A. Braunstein, R. Mulet, A. Pagnani, M. Weigt, and R. Zecchina, *Phys. Rev. E* **68**, 036702 (2003).
- [27] O.C. Martin, M. Mézard, and O. Rivoire, *Phys. Rev. Lett.* **93**, 217205 (2004).
- [28] A. Montanari, *Eur. Phys. J. B* **B23**, 121 (2001).
- [29] M. Mézard, F. Ricci-Tersenghi, and R. Zecchina, *J. Stat. Phys.* **111**, 505 (2003).
- [30] S. Cocco, O. Dubois, J. Mandler, and R. Monasson, *Phys. Rev. Lett.* **90**, 047205 (2003).
- [31] M. Mézard, T. Mora, and R. Zecchina (to be published).

“Pairs of SAT assignments in Random Boolean Formulæ”

cond-mat/0506053

Avant-propos. Dans une version soumise sur le serveur de preprints arXiv, ce papier comportait ce que nous pensions être une preuve correcte du seuil abrupt de la x -satisfaisabilité (anciennement Théorème 2, devenu Conjecture 4). Cette preuve s’avère erronée, et nous travaillons actuellement à la réparer. La version qui suit corrige le preprint, en présentant le seuil abrupt comme une conjecture.

Pairs of SAT Assignment in Random Boolean Formulæ

Thierry Mora

*Laboratoire de Physique Théorique et Modèles Statistiques, bâtiment 100,
Université Paris-Sud, F-91405 Orsay, France.*

Marc Mézard

*Laboratoire de Physique Théorique et Modèles Statistiques, bâtiment 100,
Université Paris-Sud, F-91405 Orsay, France.*

Riccardo Zecchina

*Abdus Salam International Center for Theoretical Physics, Strada Costiera 11,
34100 Trieste, Italy*

Abstract

We investigate geometrical properties of the random K -satisfiability problem using the notion of x -satisfiability: a formula is x -satisfiable if there exist two SAT assignments differing in Nx variables. For large enough K , we prove that there exists a region of clause density, below the satisfiability threshold, where the landscape of Hamming distances between SAT assignments experiences a gap: pairs of SAT-assignments exist at small x , and around $x = \frac{1}{2}$ with finite probability, but they do not exist at intermediate values of x . This result is consistent with the clustering scenario which is at the heart of the recent heuristic analysis of satisfiability using statistical physics analysis (the cavity method), and its algorithmic counterpart (the survey propagation algorithm). Our method uses elementary probabilistic arguments (first and second moment methods), and might be useful in other problems of computational and physical interest where similar phenomena appear.

Key words: satisfiability, clustering

PACS: 75.10.Nr, 75.40.-s, 75.40.Mg

1 Introduction and outline

Consider a string of Boolean variables — or equivalently a string of *spins* — of size N : $\vec{\sigma} = \{\sigma_i\} \in \{-1, 1\}^N$. Call a K -clause a disjunction binding K of these

Boolean variables in such a way that one of their 2^K joint assignments is set to FALSE, and all the others to TRUE. A formula in a conjunctive normal form (CNF) is a conjunction of such clauses. The satisfiability problem is stated as: does there exist a truth assignment $\vec{\sigma}$ that satisfies this formula? A CNF formula is said to be *satisfiable* (SAT) if this is the case, and *unsatisfiable* (UNSAT) otherwise.

The satisfiability problem is often viewed as the canonical constraint satisfaction problem (CSP). It is the first problem to have been shown NP-complete [5], i.e. at least as hard as any problem for which a solution can be checked in polynomial time.

The $P \neq NP$ conjecture states that no general polynomial-time algorithm exists that can decide whether a formula is SAT or UNSAT. However formulas which are encountered in practice can often be solved easily. In order to understand properties of some typical families of formulas, one introduces a probability measure on the set of instances. In the random K -SAT problem, one generates a random K -CNF formula $F_K(N, M)$ as a conjunction of $M = N\alpha$ K -clauses, each of them being uniformly drawn from the $2^K \binom{N}{K}$ possibilities. In the recent years the random K -satisfiability problem has attracted much interest in computer science and in statistical physics. Its most striking feature is certainly its sharp threshold.

Throughout this paper, ‘with high probability’ (w.h.p.) means with a probability which goes to one as $N \rightarrow \infty$.

Conjecture 1 (Satisfiability Threshold Conjecture) *For all $K \geq 2$, there exists $\alpha_c(K)$ such that:*

- *if $\alpha < \alpha_c(K)$, $F_K(N, N\alpha)$ is satisfiable w.h.p.*
- *if $\alpha > \alpha_c(K)$, $F_K(N, N\alpha)$ is unsatisfiable w.h.p.*

The random K -SAT problem, for N large and α close to $\alpha_c(K)$, provides instances of very hard CNF formulas that can be used as benchmarks for algorithms. For such hard ensembles, the study of the typical complexity could be crucial for the understanding of the usual ‘worst-case’ complexity.

Although Conjecture 1 remains unproved, Friedgut established the existence of a non-uniform sharp threshold [11].

Theorem 1 (Friedgut) *For each $K \geq 2$, there exists a sequence $\alpha_N(K)$ such that for all $\epsilon > 0$:*

$$\lim_{N \rightarrow \infty} \mathbf{P}(F_K(N, N\alpha) \text{ is satisfiable}) = \begin{cases} 1 & \text{if } \alpha = (1 - \epsilon)\alpha_N(K) \\ 0 & \text{if } \alpha = (1 + \epsilon)\alpha_N(K). \end{cases} \quad (1)$$

A lot of efforts have been devoted to finding tight bounds for the threshold. The best upper bounds so far were derived using first moment methods [12, 13], and the best lower bounds were obtained by second moment methods [16, 17]. Using these bounds, it was shown that $\alpha_c(K) = 2^K \ln(2) - O(K)$ as $K \rightarrow \infty$.

On the other hand, powerful, self-consistent, but non-rigorous tools from statistical physics were used to predict specific values of $\alpha_c(K)$, as well as heuristical asymptotic expansions for large K [19, 20, 21]. The *cavity method* [18], which provides these results, relies on several unproven assumptions motivated by spin-glass theory, the most important of which is the partition of the space of SAT-assignments into many *states* or *clusters* far away from each other (with Hamming distance greater than cN as $N \rightarrow \infty$), in the so-called hard-SAT phase.

So far, the existence of such a clustering phase has been shown rigorously in the simpler case of the random XORSAT problem [32, 31, 33] in compliance with the prediction of the cavity method, but its existence is predicted in many other problems, such as q -colorability [26, 27] or the Multi-Index Matching Problem [28]. At the heuristic level, clustering is an important phenomenon, often held responsible for entrapping local search algorithm into non-optimal metastable states [25]. It is also a limiting feature for the belief propagation iterative decoding algorithms in Low Density Parity Check Codes [29, 30].

In this paper we provide a rigorous analysis of some geometrical properties of the space of SAT-assignments in the random K -SAT problem. This study complements the results of [34], and its results are consistent with the clustering scenario. A new characterizing feature of CNF formulas, the ‘ x -satisfiability’, is proposed, which carries information about the spectrum of distances between SAT-assignments. The x -satisfiability property is studied thoroughly using first and second moment methods previously developed for the satisfiability threshold.

The Hamming distance between two assignments $(\vec{\sigma}, \vec{\tau})$ is defined by

$$d_{\vec{\sigma}\vec{\tau}} = \frac{N}{2} - \frac{1}{2} \sum_{i=1}^N \sigma_i \tau_i . \quad (2)$$

(Throughout the paper the term ‘distance’ will always refer to the Hamming distance.) Given a random formula $F_K(N, N\alpha)$, we define a ‘SAT- x -pair’ as a pair of assignments $(\vec{\sigma}, \vec{\tau}) \in \{-1, 1\}^{2N}$, which both satisfy F , and which are at a fixed distance specified by x as follows:

$$d_{\vec{\sigma}\vec{\tau}} \in [Nx - \epsilon(N), Nx + \epsilon(N)]. \quad (3)$$

Here x is the proportion of distinct values between the two configurations, which we keep fixed as N and d go to infinity. The resolution $\epsilon(N)$ has to be ≥ 1

and sub-extensive: $\lim_{N \rightarrow \infty} \epsilon(N)/N = 0$, but its precise form is unimportant for our large N analysis. For example we can choose $\epsilon(N) = \sqrt{N}$.

Definition 1 *A CNF formula is x -satisfiable if it possesses a SAT- x -pair.*

Note that for $x = 0$, x -satisfiability is equivalent to satisfiability, while for $x = 1$, it is equivalent to Not-All-Equal satisfiability, where each clause must contain at least one satisfied literal and at least one unsatisfied literal [16].

The clustering property found heuristically in [20, 19] suggests the following:

Conjecture 2 *For all $K \geq K_0$, there exist $\alpha_1(K)$, $\alpha_2(K)$, with $\alpha_1(K) < \alpha_2(K)$, such that: for all $\alpha \in (\alpha_1(K), \alpha_2(K))$, there exist $x_1(K, \alpha) < x_2(K, \alpha) < x_3(K, \alpha)$ such that:*

- *for all $x \in [0, x_1(K, \alpha)] \cup [x_2(K, \alpha), x_3(K, \alpha)]$, a random formula $F_K(N, N\alpha)$ is x -satisfiable w.h.p.*
- *for all $x \in [x_1(K, \alpha), x_2(K, \alpha)] \cup [x_3(K, \alpha), 1]$, a random formula $F_K(N, N\alpha)$ is x -unsatisfiable w.h.p.*

Let us give a geometrical interpretation of this conjecture. The space of SAT-assignments is partitioned into non-empty regions whose diameter is smaller than x_1 ; the distance between any two of these regions is at least x_2 , while x_3 is the maximum distance between any pair of SAT-assignments. This interpretation is compatible with the notion of clusters used in the statistical physics approach. It should also be mentioned that in a contribution posterior to this work [35], the number of regions was shown to be exponential in the size of the problem, further supporting the statistical mechanics picture.

Conjecture 2 can be rephrased in a slightly different way, which decomposes it into two steps. The first step is to state the *Satisfiability Threshold Conjecture* for pairs:

Conjecture 3 *For all $K \geq 2$ and for all x , $0 < x < 1$, there exists an $\alpha_c(K, x)$ such that:*

- *if $\alpha < \alpha_c(x)$, $F_K(N, N\alpha)$ is x -satisfiable w.h.p.*
- *if $\alpha > \alpha_c(x)$, $F_K(N, N\alpha)$ is x -unsatisfiable w.h.p.*

The second step conjectures that for K large enough, as a function of x , the function $\alpha_c(K, x)$ is non monotonic and has two maxima: a local maximum at a value $x_M(K) < 1$, and a global maximum at $x = 0$.

For our purpose Conjecture 3 can be weakened by only supposing the existence of a non-uniform threshold:

Conjecture 4 *For each $K \geq 2$ and x , $0 < x < 1$, there exists a sequence*

$\alpha_N(K, x)$ such that for all $\epsilon > 0$:

$$\lim_{N \rightarrow \infty} \mathbf{P}(F_K(N, N\alpha) \text{ is } x\text{-satisfiable}) = \begin{cases} 1 & \text{if } \alpha = (1 - \epsilon)\alpha_N(K, x) \\ 0 & \text{if } \alpha = (1 + \epsilon)\alpha_N(K, x). \end{cases} \quad (4)$$

In this paper we obtain two functions, $\alpha_{LB}(K, x)$ and $\alpha_{UB}(K, x)$, such that:

- For $\alpha > \alpha_{UB}(K, x)$, a random K -CNF $F_K(N, N\alpha)$ is x -unsatisfiable w.h.p.
- For $\alpha < \alpha_{LB}(K, x)$, a random K -CNF $F_K(N, N\alpha)$ is x -satisfiable with probability bounded away from zero.

The function $\alpha_{UB}(K, x)$ is an upper bound of $\alpha_N(K, x)$ as N tends to infinity, and $\alpha_{LB}(K, x)$ is a lower bound under Conjecture 4. Numerical computations of these bounds indicate that $\alpha_N(K, x)$ is non monotonic as a function of x for $K \geq 8$, as illustrated in Fig. 1. More precisely, we prove

Theorem 2 *For all $\epsilon > 0$, there exists K_0 such that for all $K \geq K_0$,*

$$\min_{x \in (0, \frac{1}{2})} \alpha_{UB}(K, x) \leq (1 + \epsilon) \frac{2^K \ln 2}{2}, \quad (5)$$

$$\alpha_{LB}(K, 0) \geq (1 - \epsilon) 2^K \ln 2, \quad (6)$$

$$\alpha_{LB}(K, 1/2) \geq (1 - \epsilon) 2^K \ln 2. \quad (7)$$

This in turn shows that, for K large enough and in some well chosen interval of α below the satisfiability threshold $\alpha_c \sim 2^K \ln 2$, SAT- x -pairs exist for x close to zero w.h.p. and for $x = \frac{1}{2}$ with probability bounded away from 0, but they do not exist in the intermediate x zone. Note that Eq. (6) was established by [17].

In section 2 we establish rigorous and explicit upper bounds using the first-moment method. The existence of a gap interval is proven in a certain range of α , and bounds on this interval are found, which imply Eq. (5) in Theorem 2. Section 3 derives the lower bound, using a weighted second-moment method, as developed recently in [16, 17], and presents numerical results. In section 4 we discuss the behavior of the lower bound for large K . The case of $x = \frac{1}{2}$ is treated rigorously, and Eq. (7) in Theorem 2 is proven. Other values of x are treated at the heuristic level. We discuss our results in section 5.

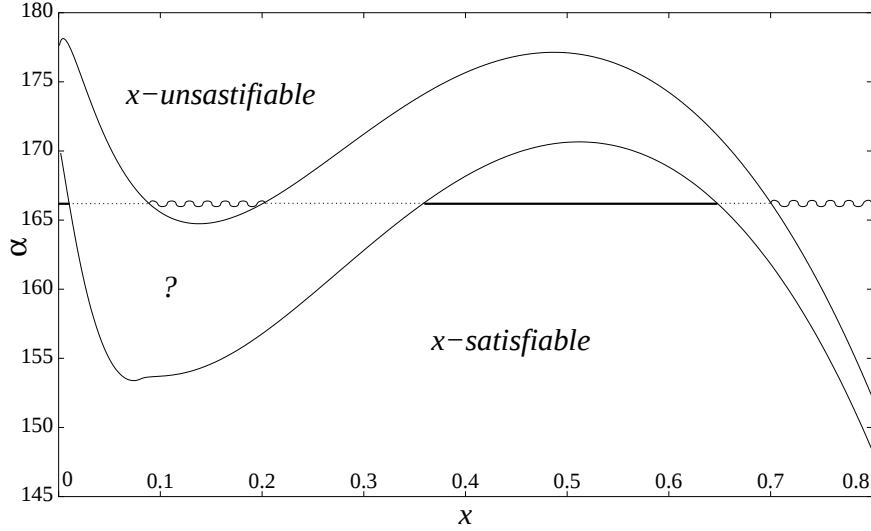


Fig. 1. Lower and Upper Bounds for $\alpha_N(K = 8, x)$. The Upper Bound is obtained by the first moment method. Above this curve there exists no SAT- x -pair, w.h.p. The Lower Bound is obtained by the second moment method. Below this curve the probability that there exist a SAT- x -pair is bounded away from 0. For $164.735 < \alpha < 170.657$, these curves confirm the existence of a clustering phase, illustrated here for $\alpha = 166.1$: solid lines represent x -sat zones, and wavy lines x -unsat zones. The x -sat zone near 0 corresponds to SAT-assignments belonging to the same region, whereas the x -sat zone around $\frac{1}{2}$ corresponds to SAT-assignments belonging to different regions. The x -unsat zone around .13 corresponds to the inter-regional gap. We recall that the best refined lower and upper bounds for the satisfiability threshold $\alpha_c(K = 8)$ from [13, 17] are respectively 173.253 and 176.596. The cavity prediction is $\alpha_c(K = 8) = 176.543$ [21].

2 Upper bound: the first moment method

The first moment method relies on Markov's inequality:

Lemma 1 *Let X be a non-negative random variable. Then*

$$\mathbf{P}(X \geq 1) \leq \mathbf{E}(X) . \quad (8)$$

We take X to be the number of pairs of SAT-assignments at fixed distance:

$$Z(x, F) = \sum_{\vec{\sigma}, \vec{\tau}} \delta(d_{\vec{\sigma}\vec{\tau}} \in [Nx + \epsilon(N), Nx - \epsilon(N)]) \delta[\vec{\sigma}, \vec{\tau} \in S(F)] , \quad (9)$$

where $F = F_K(N, N\alpha)$ is a random K -CNF formula, and $S(F)$ is the set of SAT-assignments to this formula. Throughout this paper $\delta(A)$ is an indicator function, equal to 1 if the statement A is true, equal to 0 otherwise. The expectation $\mathbf{E}$ is over the set of random K -CNF formulas. Since $Z(x, F) \geq 1$ is equivalent to ' F is x -satisfiable', (8) gives an upper bound for the probability of x -satisfiability.

The expected value of the double sum can be rewritten as:

$$\mathbf{E}(Z) = 2^N \sum_{d \in [Nx + \epsilon(N), Nx - \epsilon(N)] \cap \mathbb{N}} \binom{N}{d} \mathbf{E}[\delta(\vec{\sigma}, \vec{\tau} \in S(F))]. \quad (10)$$

where $\vec{\sigma}$ and $\vec{\tau}$ are any two assignments with Hamming distance d . We have $\delta(\vec{\sigma}, \vec{\tau} \in S(F)) = \prod_c \delta(\vec{\sigma}, \vec{\tau} \in S(c))$, where c denotes one of the M clauses. All clauses are drawn independently, so that we have:

$$\mathbf{E}(Z) \leq (2\epsilon(N) + 1)2^N \max_{d \in [Nx + \epsilon(N), Nx - \epsilon(N)] \cap \mathbb{N}} \left\{ \binom{N}{d} (\mathbf{E}[\delta(\vec{\sigma}, \vec{\tau} \in S(c))])^M \right\}, \quad (11)$$

where we have bounded the sum by the maximal term times the number of terms. $\mathbf{E}[\delta(\vec{\sigma}, \vec{\tau} \in S(c))]$ can easily be calculated and its value is: $1 - 2^{1-K} + 2^{-K}(1-x)^K + o(1)$. Indeed there are only two realizations of the clause among 2^K that do not satisfy c unless the two configurations overlap exactly on the domain of c .

Considering the normalized logarithm of this quantity,

$$F(x, \alpha) = \lim_{N \rightarrow \infty} \frac{1}{N} \ln \mathbf{E}(Z) = \ln 2 + H_2(x) + \alpha \ln (1 - 2^{1-K} + 2^{-K}(1-x)^K), \quad (12)$$

where $H_2(x) = -x \ln x - (1-x) \ln(1-x)$ is the two-state entropy function, one can deduce an upper bound for $\alpha_N(K, x)$. Indeed, $F(x, \alpha) < 0$ implies $\lim_{N \rightarrow \infty} \mathbf{P}(Z(x, F) \geq 1) = 0$. Therefore:

Theorem 3 *For each K and $0 < x < 1$, and for all α such that*

$$\alpha > \alpha_{UB}(K, x) = -\frac{\ln 2 + H_2(x)}{\ln(1 - 2^{1-K} + 2^{-K}(1-x)^K)}, \quad (13)$$

a random formula $F_K(N, N\alpha)$ is x -unsatisfiable w.h.p.

We observe numerically that a ‘gap’ (x_1, x_2 and α such that $x_1 < x < x_2 \implies F(x, \alpha) < 0$) appears for $K \geq 6$. More generally, the following results holds, which implies Eq. (5) in Theorem 2:

Theorem 4 *Let $\epsilon \in (0, 1)$, and $\{y_K\}_{K \in \mathbb{N}}$ be a sequence verifying $Ky_K \rightarrow \infty$ and $y_K = o(1)$. Denote by $H_2^{-1}(u)$ the smallest root to $H_2(x) = u$, with $u \in [0, \ln 2]$.*

There exists K_0 such that for all $K \geq K_0$, $\alpha \in [(1 + \epsilon)2^{K-1} \ln 2, \alpha_N(K))$ and $x \in [y_K, H_2^{-1}(\alpha 2^{1-K} - \ln 2 - \epsilon)] \cup [1 - H_2^{-1}(\alpha 2^{1-K} - \ln 2 - \epsilon), 1]$, $F_K(N, N\alpha)$ is x -unsatisfiable w.h.p.

Proof. Clearly $(1 + \epsilon)2^{K-1} \ln 2 < \alpha_N(K)$ since $\alpha_N(K) = 2^K \ln 2 - O_K(K)$ [17]. Observe that $(1 - y_K)^K = o(1)$. Then for all $\delta > 0$, there exists K_1 such

that for all $K \geq K_1$, $x > y_K$:

$$\alpha_{UB}(x) < (1 + \delta)2^{K-1}(\ln 2 + H_2(x)). \quad (14)$$

Inverting this inequality yields the theorem. $\square$

The choice (9) of X , although it is the simplest one, is not optimal. The first moment method only requires the condition $X \geq 1$ to be equivalent to the x -satisfiability, and better choices of X exist which allow to improve the bound. Techniques similar to the one introduced separately by Dubois and Boufkhad [13] on the one hand, and Kirousis, Kranakis and Krizanc [12] on the other hand, can be used to obtain two tighter bounds. Quantitatively, it turns out that these more elaborate bounds provide only very little improvement over the simple bound (13) (see Fig. 2). For the sake of completeness, we give without proof the simplest of these bounds:

Theorem 5 *The unique positive solution of the equation*

$$\begin{aligned} & H_2(x) + \alpha \ln \left(1 - 2^{1-K} + 2^{-K}(1-x)^K \right) \\ & + (1-x) \ln \left[2 - \exp \left(-K\alpha \frac{2^{1-K} - 2^{-K}(1-x)^{K-1}}{1 - 2^{1-K} + 2^{-K}(1-x)^K} \right) \right] \\ & + x \ln \left[2 - \exp \left(-K\alpha \frac{2^{1-K} - 2^{1-K}(1-x)^{K-1}}{1 - 2^{1-K} + 2^{-K}(1-x)^K} \right) \right] = 0 \end{aligned} \quad (15)$$

is an upper bound for $\alpha_N(K, x)$. For $x = 0$ we recover the expression of [12].

The proof closely follows that of [12] and presents no notable difficulty. We also derived a tighter bound based on the technique used in [13], gaining only a small improvement over the bound of Theorem 5 (less than .001%).

3 Lower bound: the second moment method

The second moment method uses the following consequence of Chebyshev's inequality:

Lemma 2 *If X is a non-negative random variable, one has:*

$$\mathbf{P}(X > 0) \geq \frac{\mathbf{E}(X)^2}{\mathbf{E}(X^2)}. \quad (16)$$

It is well known that the simplest choice of X as the number of SAT-assignments (in our case the number of SAT- x -pairs) is bound to fail. The intuitive reason [16, 17] is that this naive choice favors pairs of SAT-assignments

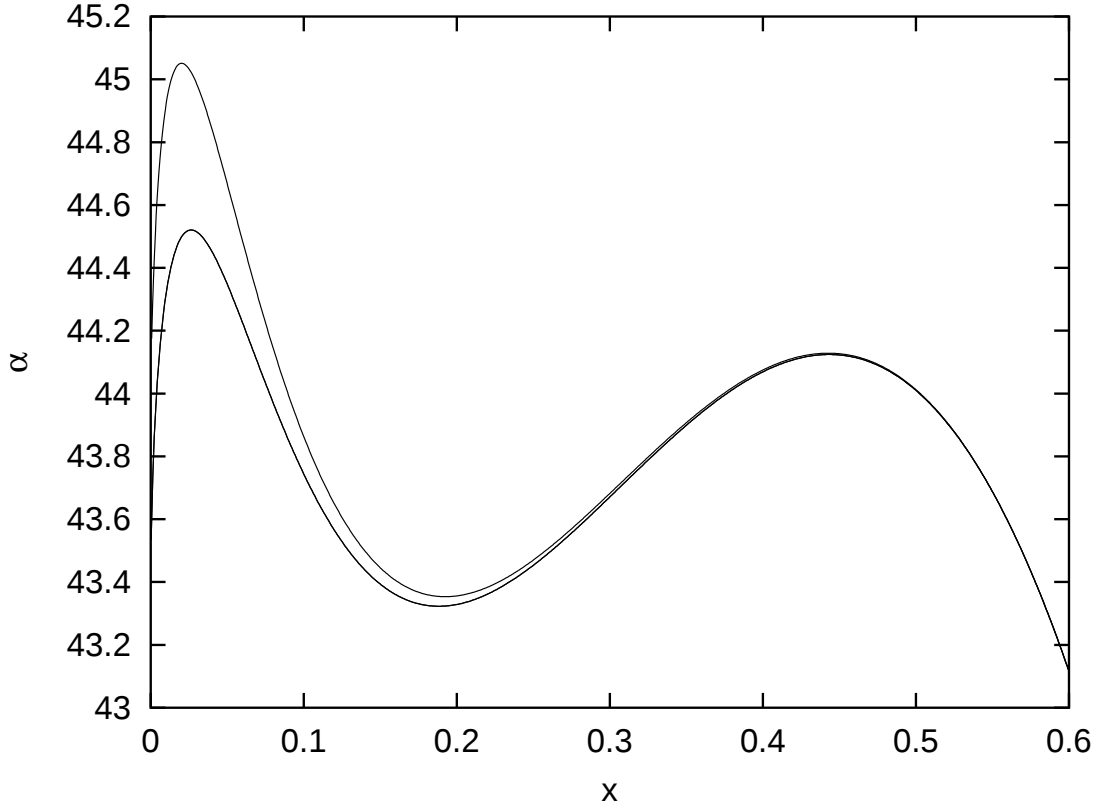


Fig. 2. Comparison between the simple upper bound (13) for $\alpha_N(K = 6, x)$ (top curve) and the refined one (bottom curve), as defined in Theorem 5.

with a great number of satisfying literals. It turns out that such assignments are highly correlated, since they tend to agree with each other, and this causes the failure of the second-moment method. In order to deal with *balanced* (with approximately half of literals satisfied) and uncorrelated pairs of assignments, one must consider a weighted sum of all SAT-assignments. Following [16, 17], we define:

$$Z(x, F) = \sum_{\vec{\sigma}, \vec{\tau}} \delta(d_{\vec{\sigma}\vec{\tau}} = \lfloor Nx \rfloor) W(\vec{\sigma}, \vec{\tau}, F), \quad (17)$$

where $\lfloor Nx \rfloor$ denotes the integer part of Nx . Note that the condition $d_{\vec{\sigma}\vec{\tau}} = \lfloor Nx \rfloor$ is stronger than Eq. (3). The weights $W(\vec{\sigma}, \vec{\tau}, F)$ are decomposed according to each clause:

$$W(\vec{\sigma}, \vec{\tau}, F) = \prod_c W(\vec{\sigma}, \vec{\tau}, c), \quad (18)$$

$$\text{with } W(\vec{\sigma}, \vec{\tau}, c) = W(\vec{u}, \vec{v}), \quad (19)$$

where $\vec{u}, \vec{v}$ are K -component vectors such that: $u_i = 1$ if the i^{th} literal of c is satisfied under $\vec{\sigma}$, and $u_i = -1$ otherwise (here we assume that the variables connected to c are arbitrarily ordered). $\vec{v}$ is defined in the same way with

respect to $\vec{\tau}$. In order to have the equivalence between $Z > 0$ and the existence of pairs of SAT-assignments, we impose the following condition on the weights:

$$W(\vec{u}, \vec{v}) = \begin{cases} 0 & \text{if } \vec{u} = (-1, \dots, -1) \text{ or } \vec{v} = (-1, \dots, -1), \\ > 0 & \text{otherwise.} \end{cases} \quad (20)$$

Let us now compute the first and second moments of Z :

Fact 1

$$\mathbf{E}(Z) = 2^N \binom{N}{\lfloor Nx \rfloor} f_1(x)^M, \quad (21)$$

where

$$f_1(x) = \mathbf{E}[W(\vec{\sigma}, \vec{\tau}, c)] \quad (22)$$

$$= 2^{-K} \sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) (1-x)^{|\vec{u} \cdot \vec{v}|} x^{K-|\vec{u} \cdot \vec{v}|}. \quad (23)$$

Here $|\vec{u}|$ is the number of indices i such that $u_i = +1$, and $\vec{u} \cdot \vec{v}$ denotes the vector $(u_1 v_1, \dots, u_K v_K)$.

Writing the second moment is a little more cumbersome:

Fact 2

$$\mathbf{E}(Z^2) = 2^N \sum_{\mathbf{a} \in V_N \cap \{0, 1/N, 2/N, \dots, 1\}^8} \frac{N!}{\prod_{i=0}^7 (N a_i)!} f_2(\mathbf{a})^M, \quad (24)$$

where

$$\begin{aligned} f_2(\mathbf{a}) &= \mathbf{E}[W(\vec{\sigma}, \vec{\tau}, c) W(\vec{\sigma}, \vec{\tau}, c)] \\ &= 2^{-K} \sum_{\vec{u}, \vec{v}, \vec{u}', \vec{v}'} W(\vec{u}, \vec{v}) W(\vec{u}', \vec{v}') \prod_{i=1}^K a_0^{\delta(u_i=v_i=u'_i=v'_i)} a_1^{\delta(u_i=v_i=u'_i \neq v'_i)} \\ &\quad a_2^{\delta(u_i=v_i=v'_i \neq u'_i)} a_3^{\delta((u_i=v_i) \neq (u'_i=v'_i))} a_4^{\delta(u_i=u'_i=v'_i \neq v_i)} \\ &\quad a_5^{\delta((u_i=u'_i) \neq (v_i=v'_i))} a_6^{\delta((u_i=v'_i) \neq (u'_i=v_i))} a_7^{\delta(u'_i=v'_i=u_i \neq u_i)} \end{aligned} \quad (25)$$

$\mathbf{a}$ is a 8-component vector giving the proportion of each type of quadruplets $(\tau_i, \sigma_i, \tau'_i, \sigma'_i)$ — $\vec{\tau}$ being arbitrarily (but without losing generality) fixed to $(1, \dots, 1)$ — as described in the following table:

	a_0	a_1	a_2	a_3	a_4	a_5	a_6	a_7
τ_i	+	+	+	+	+	+	+	+
σ_i	+	+	+	+	-	-	-	-
τ'_i	+	+	-	-	+	+	-	-
σ'_i	+	-	+	-	+	-	+	-

The set $V_N \subset [0, 1]^8$ is a simplex specified by:

$$\begin{cases} \lfloor N(a_4 + a_5 + a_6 + a_7) \rfloor = \lfloor Nx \rfloor \\ \lfloor N(a_1 + a_2 + a_5 + a_6) \rfloor = \lfloor Nx \rfloor \\ \sum_{i=0}^7 a_i = 1 \end{cases} \quad (26)$$

These three conditions (26) correspond to the normalization of the proportions and to the enforcement of the conditions $d_{\vec{\sigma}\vec{\tau}} = \lfloor Nx \rfloor$, $d_{\vec{\sigma}'\vec{\tau}'} = \lfloor Nx \rfloor$. When $N \rightarrow \infty$, $V = \bigcap_{N \in \mathbb{N}} V_N$ defines a five-dimensional simplex described by the three hyperplanes:

$$\begin{cases} a_4 + a_5 + a_6 + a_7 = x \\ a_1 + a_2 + a_5 + a_6 = x \\ \sum_{i=0}^7 a_i = 1 \end{cases} \quad (27)$$

In order to yield an asymptotic estimate of $\mathbf{E}(Z^2)$ we first use the following lemma, which results from a simple approximation of integrals by sums:

Lemma 3 *Let $\psi(\mathbf{a})$ be a real, positive, continuous function of $\mathbf{a}$, and let V_N , V be defined as previously. Then there exists a constant C_0 depending on x such that for sufficiently large N :*

$$\sum_{\mathbf{a} \in V_N \cap \{1/N, 2/N, \dots, 1\}^8} \frac{N!}{\prod_{i=0}^7 (Na_i)!} \psi(\mathbf{a})^N \leq C_0 N^{3/2} \int_V d\mathbf{a} e^{N[H_8(\mathbf{a}) + \ln \psi(\mathbf{a})]}, \quad (28)$$

where $H_8(\mathbf{a}) = -\sum_{i=1}^8 a_i \ln a_i$.

A standard Laplace method used on Eq. (28) with $\psi = 2(f_2)^\alpha$ yields:

Fact 3 *For each K, x , define:*

$$\Phi(\mathbf{a}) = H_8(\mathbf{a}) - \ln 2 - 2H_2(x) + \alpha \ln f_2(\mathbf{a}) - 2\alpha \ln f_1(x). \quad (29)$$

and let $\mathbf{a}_0 \in V$ be the global maximum of Φ restricted to V . Suppose that $\partial_{\mathbf{a}}^2 \Phi(\mathbf{a}_0)$ is definite negative. Then there exists a constant C_1 such that, for N

sufficiently large,

$$\frac{\mathbf{E}(Z)^2}{\mathbf{E}(Z^2)} \geq C_1 \exp(-N\Phi(\mathbf{a}_0)). \quad (30)$$

Obviously $\Phi(\mathbf{a}_0) \geq 0$ in general. In order to use Lemma 2, one must find the weights $W(\vec{u}, \vec{v})$ in such a way that $\max_{\mathbf{a} \in V} \Phi(\mathbf{a}) = 0$. We first notice that, at the particular point $\mathbf{a}^*$ where the two pairs are uncorrelated with each other,

$$a_0^* = a_3^* = \frac{(1-x)^2}{2}, \quad a_1^* = a_2^* = a_4^* = a_7^* = \frac{x(1-x)}{2}, \quad a_5^* = a_6^* = \frac{x^2}{2}, \quad (31)$$

we have the following properties:

- $H_8(\mathbf{a}^*) = \ln 2 + 2H_2(x)$,
- $\partial_{\mathbf{a}} H_8(\mathbf{a}^*) = 0$, $\partial_{\mathbf{a}}^2 H_8(\mathbf{a}^*)$ definite negative,
- $f_1(x)^2 = f_2(\mathbf{a}^*)$ and hence $\Phi(\mathbf{a}^*) = 0$.

(Note that the derivatives $\partial_{\mathbf{a}}$ are taken in the simplex V). So the weights must be chosen in such a way that $\mathbf{a}^*$ be the global maximum of Φ . A necessary condition is that $\mathbf{a}^*$ be a local maximum, which entails $\partial_{\mathbf{a}} f_2(\mathbf{a}^*) = 0$.

Using the fact that the number of common values between four vectors $\vec{u}, \vec{v}, \vec{u}', \vec{v}' \in \{-1, 1\}^K$ can be written as:

$$\frac{1}{8} \left(K + \vec{u} \cdot \vec{v} + \vec{u} \cdot \vec{u}' + \vec{u} \cdot \vec{v}' + \vec{v} \cdot \vec{u}' + \vec{v} \cdot \vec{v}' + \vec{u}' \cdot \vec{v}' + \overrightarrow{u \cdot v} \cdot \overrightarrow{u' \cdot v'} \right) \quad (32)$$

we deduce from $\partial_{\mathbf{a}} f_2(\mathbf{a}^*) = 0$ the condition:

$$\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) \begin{cases} \vec{u} \\ \vec{v} \end{cases} (1-x)^{|\overrightarrow{u \cdot v}|} x^{K-|\overrightarrow{u \cdot v}|} = 0, \quad (33)$$

$$\begin{aligned} 0 &= K(2x-1)^2 \left[\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) (1-x)^{|\overrightarrow{u \cdot v}|} x^{K-|\overrightarrow{u \cdot v}|} \right]^2 \\ &\quad + \left[\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) \overrightarrow{u \cdot v} (1-x)^{|\overrightarrow{u \cdot v}|} x^{K-|\overrightarrow{u \cdot v}|} \right]^2 \\ &\quad + 2(2x-1) \left[\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) \vec{u} \cdot \vec{v} (1-x)^{|\overrightarrow{u \cdot v}|} x^{K-|\overrightarrow{u \cdot v}|} \right] \\ &\quad \times \left[\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) (1-x)^{|\overrightarrow{u \cdot v}|} x^{K-|\overrightarrow{u \cdot v}|} \right]. \end{aligned} \quad (34)$$

If we suppose that W is invariant under simultaneous and identical permutations of the u_i or of the v_i (which we must, since the ordering of the variables by the label i is arbitrary), the K components of all vectorial quantities in Eqs. (33), (34) should be equal. Then we obtain equivalently:

$$\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) (2|\vec{u}| - K) (1 - x)^{|\vec{u}\vec{v}|} x^{K-|\vec{u}\vec{v}|} = 0 \quad \text{and} \quad \vec{u} \leftrightarrow \vec{v}, \quad (35)$$

$$\sum_{\vec{u}, \vec{v}} W(\vec{u}, \vec{v}) (K(2x - 1) + \vec{u} \cdot \vec{v}) (1 - x)^{|\vec{u}\vec{v}|} x^{K-|\vec{u}\vec{v}|} = 0, \quad (36)$$

We choose the following simple form for $W(\vec{u}, \vec{v})$:

$$W(\vec{u}, \vec{v}) = \begin{cases} 0 & \text{if } \vec{u} = (-1, \dots, -1) \quad \text{or} \quad \vec{v} = (-1, \dots, -1), \\ \lambda^{|\vec{u}|+|\vec{v}|} \nu^{|\vec{u}\vec{v}|} & \text{otherwise.} \end{cases} \quad (37)$$

Although this choice is certainly not optimal, it turns out particularly tractable. Eqs. (35) and (36) simplify to:

$$\begin{aligned} [\nu(1 - x)]^{K-1} &= (\lambda^2 + 1 - 2\lambda\nu) \left(2\lambda x + \nu(1 - x)(1 + \lambda^2) \right)^{K-1} \\ (\nu(1 - x) + \lambda x)^{K-1} &= (1 - \lambda\nu) \left(2\lambda x + \nu(1 - x)(1 + \lambda^2) \right)^{K-1}. \end{aligned} \quad (38)$$

We found numerically a unique solution $\lambda > 0, \nu > 0$ to these equations for any value of $K \geq 2$ that we checked.

Fixing (λ, ν) to a solution of (38), we seek the largest value of α such that the local maximum $\mathbf{a}^*$ is a global maximum, i.e. such that there exists no $\mathbf{a} \in V$ with $\Phi(\mathbf{a}) > 0$. To proceed one needs analytical expressions for $f_1(x)$ and $f_2(\mathbf{a})$. f_1 simply reads:

$$\begin{aligned} f_1(x) &= 2^{-K} \left((1 - x)\nu(1 + \lambda^2) + 2x\lambda \right)^K - 2 \cdot 2^{-K} (x\lambda + (1 - x)\nu)^K \\ &\quad + 2^{-K} ((1 - x)\nu)^K. \end{aligned} \quad (39)$$

f_2 is calculated by Sylvester's formula, but its expression is long and requires preliminar notations. We index the 16 possibilities for (u_i, v_i, u'_i, v'_i) by a number $r \in \{0, \dots, 15\}$ defined as:

$$r = 8 \frac{1 - u_i}{2} + 4 \frac{1 - v_i}{2} + 2 \frac{1 - u'_i}{2} + \frac{1 - v'_i}{2}. \quad (40)$$

For each index r , define

$$l(r) = \delta(u_i = 1) + \delta(v_i = 1) + \delta(u'_i = 1) + \delta(v'_i = 1), \quad (41)$$

$$n(r) = \delta(u_i v_i = 1) + \delta(u'_i v'_i = 1), \quad (42)$$

and

$$z_r = \lambda^{l(r)} \nu^{n(r)} \times \begin{cases} a_r & \text{if } r \leq 7 \\ a_{15-r} & \text{if } r \geq 8 \end{cases}. \quad (43)$$

Also define the four following subsets of $\{0, \dots, 15\}$: A_0 is the set of indices r corresponding to quadruplets of the form $(-1, v_i, u'_i, v'_i)$. $A_0 = \{r \in \{0, \dots, 15\} \mid u_i = -1\}$. Similarly, $A_1 = \{r \mid v_i = -1\}$, $A_2 = \{r \mid u'_i = -1\}$ and $A_3 = \{r \mid v'_i = -1\}$.

Then f_2 is given by:

$$\begin{aligned} 2^K f_2(\mathbf{a}) = & \left(\sum_{j=0}^{15} z_j \right)^K - \sum_{k=0}^3 \left(\sum_{j \in A_k} z_j \right)^K + \sum_{0 \leq k < k' \leq 3} \left(\sum_{j \in A_k \cap A_{k'}} z_j \right)^K \\ & - \sum_{0 \leq k < k' < k'' \leq 3} \left(\sum_{j \in A_k \cap A_{k'} \cap A_{k''}} z_j \right)^K + \left(\sum_{j \in A_0 \cap A_1 \cap A_2 \cap A_3} z_j \right)^K. \end{aligned} \quad (44)$$

We can now state our lower-bound result:

Lemma 4 *Let $\alpha_+ \in (0, +\infty]$ be the smallest α such that $\partial_{\mathbf{a}}^2 \Phi(\mathbf{a}^*)$ is not definite negative. For each K and $x \in (0, 1)$, and for all $\alpha \leq \alpha_{LB}(K, x)$, with*

$$\alpha_{LB}(K, x) = \min \left[\alpha_+, \inf_{\mathbf{a} \in V_+} \frac{\ln 2 + 2H_2(x) - H_8(\mathbf{a})}{\ln f_2(\mathbf{a}) - 2 \ln f_1(x)} \right], \quad (45)$$

where $V_+ = \{\mathbf{a} \in V \mid f_2(\mathbf{a}) > f_1^2(1/2)\}$, and where (λ, ν) is chosen to be a positive solution of (38), the probability that a random formula $F_K(N, N\alpha)$ is x -satisfiable is bounded away from 0 as $N \rightarrow \infty$.

This is a straightforward consequence of the expression (29) of $\Phi(\mathbf{a})$.

If Conjecture 4 were true, Lemma 4 would imply the x -satisfiability w.h.p for all $\alpha < \alpha_{LB}(K, x)$:

Proposition 4 *For all $\alpha < \alpha_{LB}(K, x)$ defined in Lemma 4, a random K -CNF formula $F_K(N, N\alpha)$ is x -satisfiable w.h.p., unless x -satisfiability has a coarse threshold.*

We devised several numerical strategies to evaluate $\alpha_{LB}(K, x)$. The implementation of Powell's method on each point of a grid of size $\mathcal{N}^5$ ($\mathcal{N} = 10, 15, 20$) on V turned out to be the most efficient and reliable. The results are given

by Fig. 1 for $K = 8$, the smallest K such that the picture given by Conjecture 2 is confirmed. We found a clustering phenomenon for all the values of $K \geq 8$ that we checked. In the following we shall provide a rigorous estimate of $\alpha_{LB}(K, \frac{1}{2})$ at large K .

4 Large K analysis

4.1 Asymptotics for $x = \frac{1}{2}$

The main result of this section is contained in the following theorem, which implies Eq. (7) in Theorem 2:

Theorem 6 *The large K asymptotics of $\alpha_{LB}(K, x)$ at $x = 1/2$ is given by:*

$$\alpha_{LB}(K, 1/2) \sim 2^K \ln 2. \quad (46)$$

The proof primarily relies on the following results:

Fact 5 *Let $\nu = 1$ and λ be the unique positive root of:*

$$(1 - \lambda)(1 + \lambda)^{K-1} - 1 = 0. \quad (47)$$

Then (λ, ν) is solution to (38) with $x = \frac{1}{2}$ and one has, at large K :

$$\lambda - 1 \sim -2^{1-K}. \quad (48)$$

Lemma 5 *Let $x = \frac{1}{2}$. There exist $K_0 > 0$, $C_1 > 0$ and $C_2 > 0$ such that for all $K \geq K_0$, and for all $\mathbf{a} \in V$ s.t. $|\mathbf{a} - \mathbf{a}^*| < 1/8$,*

$$|\ln f_2(\mathbf{a}) - 2 \ln f_1(1/2)| \leq K^2 C_1 |\mathbf{a} - \mathbf{a}^*|^2 2^{-2K} + C_2 |\mathbf{a} - \mathbf{a}^*|^3 2^{-K} \quad (49)$$

Lemma 6 *Let $x = \frac{1}{2}$. There exist $K_0 > 0$, $C_0 > 0$ such that for $K \geq K_0$, for all $\mathbf{a} \in V$,*

$$\begin{aligned} |\ln f_2(\mathbf{a}) - 2 \ln f_1(1/2)| \leq 2^{-K} & \left[(a_0 + a_1 + a_4 + a_5)^K + (a_0 + a_2 + a_4 + a_6)^K \right. \\ & \left. + (a_0 + a_1 + a_6 + a_7)^K + (a_0 + a_2 + a_5 + a_7)^K \right] + C_0 K 2^{-2K} \end{aligned} \quad (50)$$

The proofs of these lemmas are deferred to sections 4.3 and 4.4.

4.2 Proof of Theorem 6

We first show that $\partial_{\mathbf{a}}^2 \Phi(\mathbf{a}^*)$ is definite negative for all $\alpha < 2^K$, when K is sufficiently large. Indeed $\partial_{\mathbf{a}}^2 H_8(\mathbf{a}^*)$ is definite negative and its largest eigenvalue is -4 . Using Lemma 5, for $\mathbf{a} \in V$ close enough to $\mathbf{a}^*$:

$$\Phi(\mathbf{a}) \leq -2|\mathbf{a} - \mathbf{a}^*|^2 + \alpha C_1 |\mathbf{a} - \mathbf{a}^*|^2 K^2 2^{-2K} + \alpha C_2 |\mathbf{a} - \mathbf{a}^*|^3 2^{-K}. \quad (51)$$

Therefore

$$\Phi(\mathbf{a}) \leq -|\mathbf{a} - \mathbf{a}^*|^2 \quad \text{for } K \text{ large enough, } |\mathbf{a} - \mathbf{a}^*| < \frac{1}{2C_2} \text{ and } \alpha < 2^K. \quad (52)$$

Using Theorem 4, we need to find the minimum, for $a \in V_+$, of

$$G(K, \mathbf{a}) \equiv \frac{3 \ln 2 - H_8(\mathbf{a})}{\ln f_2(\mathbf{a}) - 2 \ln f_1(1/2)}. \quad (53)$$

We shall show that

$$\inf_{\mathbf{a} \in V_+} G(K, \mathbf{a}) \sim 2^K \ln 2. \quad (54)$$

We divide this task in two parts. The first part states that there exists $R > 0$ and K_1 such that for all $K \geq K_1$, and for all $\mathbf{a} \in V_+$ such that $|\mathbf{a} - \mathbf{a}^*| < R$, $G(K, \mathbf{a}) > 2^K$. This is a consequence of Lemma 5; using the fact that $3 \ln 2 - H_8(\mathbf{a}) \geq |\mathbf{a} - \mathbf{a}^*|^2$ for $\mathbf{a}$ close enough to $\mathbf{a}^*$, one obtains:

$$G(K, \mathbf{a}) \geq \frac{2^K}{C_1 K^2 2^{-K} + C_2 |\mathbf{a} - \mathbf{a}^*|} \quad (55)$$

which, for K large enough and $\mathbf{a}$ close enough to $\mathbf{a}^*$, is greater than 2^K .

The second part deals with the case where $\mathbf{a}$ is far from $\mathbf{a}^*$, i.e. $|\mathbf{a} - \mathbf{a}^*| > R$. First we put a bound on the numerator of $G(\mathbf{a})$: there exists a constant $C_3 > 0$ such that for all $\mathbf{a} \in V$ s.t. $|\mathbf{a} - \mathbf{a}^*| > R$, one has $3 \ln 2 - H_8(\mathbf{a}) > C_3$.

Looking at Eq. (50), it is clear that, in order to minimize $G(K, \mathbf{a})$, $\mathbf{a}$ should be ‘close’ to at least one the four hyperplanes defined by

$$\begin{aligned} a_0 + a_1 + a_4 + a_5 &= 1, & a_0 + a_2 + a_4 + a_6 &= 1, \\ a_0 + a_1 + a_6 + a_7 &= 1, & a_0 + a_2 + a_5 + a_7 &= 1. \end{aligned} \quad (56)$$

More precisely, we say for instance that $\mathbf{a}$ is *close to* the first hyperplane defined above iff

$$a_0 + a_1 + a_4 + a_5 > 1 - K^{-1/2} \quad (57)$$

Now suppose that $\mathbf{a}$ is *not* close to that hyperplane. Then the corresponding term goes to 0:

$$(a_0 + a_1 + a_4 + a_5)^K \leq (1 - K^{-1/2})^K \sim \exp(-\sqrt{K}) \quad \text{as } K \rightarrow \infty. \quad (58)$$

We classify all possible cases according to the number of hyperplanes $\mathbf{a} \in V_+$ is close to:

- $\mathbf{a}$ is close to none of the hyperplanes. Then

$$G(K, \mathbf{a}) \geq \frac{2^K C_3}{4 \exp(-\sqrt{K}) + C_0 K 2^{-K}} > 2^K \quad \text{for } K \text{ large enough.} \quad (59)$$

- $\mathbf{a}$ is close to one hyperplane only, e.g. the first hyperplane $a_0 + a_1 + a_4 + a_5 = 1$ (the other hyperplanes are treated equivalently). As $\sum_{i=0}^7 a_i = 0$, one has

$$a_2 < K^{-1/2}, \quad a_3 < K^{-1/2}, \quad a_6 < K^{-1/2}, \quad a_7 < K^{-1/2}. \quad (60)$$

This implies $H_8(\mathbf{a}) < 2 \ln 2 + 2 \ln K / \sqrt{K}$, and we get:

$$G(K, \mathbf{a}) \geq \frac{2^K [\ln 2 - 2 \ln K / \sqrt{K}]}{1 + C_0 K 2^{-K} + 3 e^{-\sqrt{K}}} \geq 2^K (\ln 2) [1 - 3 \ln K / \sqrt{K}] \quad (61)$$

for sufficiently large K .

- $\mathbf{a}$ is close to two hyperplanes. It is easy to check that these hyperplanes must be either the first and the fourth ones, or the second and the third ones. In the first case we have $a_0 + a_5 > 1 - 3/\sqrt{K}$ and in the second case $a_0 + a_6 > 1 - 3/\sqrt{K}$. Both cases imply: $H_8(\mathbf{a}) < \ln 2 + 3 \ln K / \sqrt{K}$. One thus obtains:

$$G(K, \mathbf{a}) \geq \frac{2^K [2 \ln 2 - 3 \ln K / \sqrt{K}]}{2 + C_0 K 2^{-K} + 2 e^{-\sqrt{K}}} \geq 2^K (\ln 2) [1 - 3 \ln K / \sqrt{K}]. \quad (62)$$

- One can check that $\mathbf{a}$ cannot be close to more than two hyperplanes.

To sum up, we have proved that for K large enough, for all $\mathbf{a} \in V_+$,

$$G(K, \mathbf{a}) \geq 2^K (\ln 2) [1 - 3 \ln K / \sqrt{K}], \quad (63)$$

Clearly, $\alpha_{LB}(K, 1/2) = \inf_{\mathbf{a} \in V_+} G(K, \mathbf{a}) < \alpha_{UB}(K, 1/2)$. Since from Theorem 3 we know that $\alpha_{UB}(K, 1/2) \sim 2^K \ln 2$, this proves Eq. (54).

4.3 Proof of Lemma 5

Let $x = \frac{1}{2}$ and choose $\nu = 1$ and λ the unique positive root of Eq. (47). Let $\epsilon_i = a_i - 1/8$, and $\boldsymbol{\epsilon} = (\epsilon_0, \dots, \epsilon_7)$. We expand $f_2(\mathbf{a})$ in series of $\boldsymbol{\epsilon}$. The zeroth

order term is $f_2(1/8, \dots, 1/8) = f_1^2(1/2)$. The first order term vanishes. We thus get:

$$f_2(\mathbf{a}) = f_1^2(1/2) + B_0 - B_1 + B_2 - B_3 + B_4, \quad (64)$$

with

$$B_0 = \sum_{q=2}^K \binom{K}{q} \left(\frac{1}{2} \sum_{i=0}^7 p_i(\lambda) \epsilon_i \right)^q \left[\frac{1+\lambda}{2} \right]^{4(K-q)}, \quad (65)$$

$$B_1 = 2^{-K} \sum_{a=1}^4 \sum_{q=2}^K \binom{K}{q} \left[\sum_{i=0}^7 (\lambda^{\ell_{ai}} - 1) \epsilon_i \right]^q \left[\frac{1+\lambda}{2} \right]^{3(K-q)}, \quad (66)$$

$$B_2 = 2^{-2K} \sum_{a=1}^6 \sum_{q=2}^K \binom{K}{q} [2r_a(\lambda, \epsilon)]^q \left[\frac{1+\lambda}{2} \right]^{2(K-q)}, \quad (67)$$

$$B_3 = 2^{-3K} \sum_{a=1}^4 \sum_{q=2}^K \binom{K}{q} [4s_a(\lambda, \epsilon)]^q \left[\frac{1+\lambda}{2} \right]^{K-q}, \quad (68)$$

$$B_4 = 2^{-4K} \sum_{k=2}^K (8\epsilon_0)^q. \quad (69)$$

In B_0 , $p_i(\lambda) = \lambda^{l(i)} + \lambda^{l(15-i)} - 2 - 4(\lambda - 1)$. We have used the fact that $\sum_{i=0}^7 \epsilon_i = 0$. Using $l(i) + l(15-i) = 4$, one obtains $|p_i(\lambda)| \leq 11(\lambda - 1)^2 \leq 11 \cdot 2^{4-2K}$, since $|\lambda - 1| \leq 2^{2-K}$ for K large enough, by virtue of Lemma 5.

In B_1 , we have used again $\sum_{i=0}^7 \epsilon_i = 0$. ℓ_{ai} is either $l(i)$ or $l(15-i)$, depending on a . In both cases $|\lambda^{\ell_{ai}} - 1| \leq 4|\lambda - 1| \leq 2^{4-K}$. In B_2 and B_3 , the expressions of $r_a(\lambda, \epsilon)$ and $s_a(\lambda, \epsilon)$ are given by:

$$\begin{aligned} r_1 &= \epsilon_0 + \lambda(\epsilon_1 + \epsilon_2) + \lambda^2 \epsilon_3, & r_2 &= \epsilon_0 + \lambda(\epsilon_1 + \epsilon_4) + \lambda^2 \epsilon_5, \\ r_3 &= \epsilon_0 + \lambda(\epsilon_2 + \epsilon_4) + \lambda^2 \epsilon_6, & r_4 &= \epsilon_0 + \lambda(\epsilon_1 + \epsilon_7) + \lambda^2 \epsilon_6, \\ r_5 &= \epsilon_0 + \lambda(\epsilon_2 + \epsilon_7) + \lambda^2 \epsilon_5, & r_6 &= \epsilon_0 + \lambda(\epsilon_4 + \epsilon_7) + \lambda^2 \epsilon_3, \end{aligned} \quad (70)$$

$$s_1 = \epsilon_0 + \lambda \epsilon_1, \quad s_2 = \epsilon_0 + \lambda \epsilon_2, \quad s_3 = \epsilon_0 + \lambda \epsilon_4, \quad s_4 = \epsilon_0 + \lambda \epsilon_7. \quad (71)$$

In order to prove Lemma 5 we will use the following fact:

Fact 6 *Let y be a real variable such that $|y| \leq 1$. Then*

$$\left| \sum_{k=2}^K \binom{K}{k} y^k \right| \leq \frac{K(K-1)}{2} y^2 + 2^K |y|^3. \quad (72)$$

One has $|2r_a| \leq 8|\epsilon|$, $|4s_a| \leq 8|\epsilon|$, and $|8\epsilon_0| \leq 8|\epsilon|$. Therefore, for $|\epsilon| < 1/8$, one can write:

$$|B_0| \leq \frac{K(K-1)}{2} (11 \cdot 2^6)^2 2^{-4K} |\epsilon|^2 + (11 \cdot 2^6)^3 2^{-5K} |\epsilon|^3 \quad (73)$$

$$|B_1| \leq 4 \frac{K(K-1)}{2} 2^{14} 2^{-3K} |\epsilon|^2 + 2^{21} 2^{-3K} |\epsilon|^3 \quad (74)$$

$$|B_i| \leq \binom{4}{i} \frac{K(K-1)}{2} 2^6 2^{-iK} |\epsilon|^2 + 2^9 2^{-(i-1)K} |\epsilon|^3 \quad \text{for } 2 \leq i \leq 4. \quad (75)$$

Observe that

$$f_1(1/2) = \left[\left(\frac{1+\lambda}{2} \right)^K - 2^{-K} \right]^2 = 1 + O(K2^{-K}) \quad (76)$$

and that for K large enough,

$$\left| \ln \frac{f_2(\mathbf{a})}{f_1^2(1/2)} \right| \leq \frac{2}{f_1(1/2)^2} \sum_{i=0}^4 |B_i|, \quad (77)$$

which proves Lemma 5.

4.4 Proof of Lemma 6

Note that the bounds on B_0 and B_1 (73), (74) remain valid for any ϵ . Therefore $B_0 = O(2^{-2K})$ and $B_1 = O(2^{-2K})$ uniformly. We bound B_3 by observing that:

$$\begin{aligned} B_3 = & 2^{-K} \left[(a_0 + \lambda a_1)^K + (a_0 + \lambda a_2)^K + (a_0 + \lambda a_4)^K + (a_0 + \lambda a_7)^K \right] \\ & - 2^{-3K} \sum_{a=1}^4 \left[\frac{1+\lambda}{2} \right]^K \left[1 + K \left(\frac{8s_a(\lambda, \epsilon)}{1+\lambda} \right) \right]. \end{aligned} \quad (78)$$

Since $(a_0 + \lambda a_1) \leq a_0 + a_1 \leq 1/2$ and likewise for the three other terms, one has $B_3 = O(2^{-2K})$ uniformly in $\mathbf{a}$. A similar argument yields $B_4 = O(2^{-2K})$. There remains B_2 , which we write as:

$$\begin{aligned} B_2 = & 2^{-K} \sum_{0 \leq k < k' \leq 3} \left(\sum_{j \in A_k \cap A_{k'}} z_j \right)^K \\ & - 2^{-2K} \sum_{a=1}^6 \left[\frac{1+\lambda}{2} \right]^{2K} \left[1 + K \left(\frac{8r_a(\lambda, \epsilon)}{(1+\lambda)^2} \right) \right] \end{aligned} \quad (79)$$

The second term of the sum is $O(K2^{-2K})$. The first term is made of six contributions. Two of them, namely $2^{-K}(a_0 + \lambda(a_1 + a_2) + \lambda^2 a_3)$ and $2^{-K}(a_0 + \lambda(a_4 + a_7) + \lambda^2 a_3)$, are $O(2^{-2K})$, because of the condition on distances. Among the four remaining contributions, we show how to deal with one of them, the

others being handled similarly. This contribution can be written as:

$$(a_0 + \lambda(a_1 + a_4) + \lambda^2 a_5)^K = (a_0 + a_1 + a_4 + a_5)^K \left(1 + \frac{(\lambda - 1)(a_1 + a_4) + (\lambda^2 - 1)a_5}{a_0 + a_1 + a_4 + a_5} \right)^K. \quad (80)$$

We distinguish two cases. Either $a_0 + a_1 + a_4 + a_5 \leq 1/2$, and we get trivially:

$$(a_0 + \lambda(a_1 + a_4) + \lambda^2 a_5)^K - (a_0 + a_1 + a_4 + a_5)^K = O(2^{-K}), \quad (81)$$

since both terms are $O(2^{-K})$; or $a_0 + a_1 + a_4 + a_5 \geq 1/2$, and then:

$$\begin{aligned} & \left| (a_0 + \lambda(a_1 + a_4) + \lambda^2 a_5)^K - (a_0 + a_1 + a_4 + a_5)^K \right| \leq \\ & \left| \left(1 + \frac{(\lambda - 1)(a_1 + a_4) + (\lambda^2 - 1)a_5}{a_0 + a_1 + a_4 + a_5} \right)^K - 1 \right| = O(K 2^{-K}). \end{aligned} \quad (82)$$

Using again Eq. (76) finishes the proof of Lemma 6. $\square$

4.5 Heuristics for arbitrary x

For arbitrary x , the function to minimize in (45) is hard to study analytically. Here we present what we believe to be the correct asymptotic expansion of $\alpha_{LB}(K, x)$ at large K . Hopefully this tentative analysis could be used as a starting point towards a rigorous analytical treatment for any x .

A careful look at the numerics suggests the following Ansatz on the position of the global maximum, at large K :

$$\begin{aligned} a_0 &= 1 - x + o(1), & a_6 &= x + o(1) \\ a_i &= o(1) & \text{for } i \neq 0, 6. \end{aligned} \quad (83)$$

A second, symmetric, maximum also exists around $a_0 = 1 - x$, $a_5 = x$. Plugging this locus into Eq. (45) leads to the following conjecture:

Conjecture 5 *For all $x \in (0, 1]$, the asymptotics of $\alpha_{LB}(x)$ is given by:*

$$\lim_{K \rightarrow \infty} 2^{-K} \alpha_{LB}(K, x) = \frac{\ln 2 + H(x)}{2}, \quad (84)$$

and the limit is uniform on any closed sub-interval of $(0, 1]$.

This conjecture is consistent with both our numerical simulations and our result at $x = \frac{1}{2}$.

5 Discussion and Conclusion

We have developed a simple and rigorous probabilistic method which paves the way towards a complete characterization of the clustered hard-SAT phase in the random satisfiability problem. Our result is consistent with the clustering picture and supports the validity of the one-step replica symmetry breaking scheme of the cavity method for $K \geq 8$.

The study of x -satisfiability has the advantage that it does not rely on a precise definition of clusters. Indeed, it is important to stress that the “appropriate” definition for clusters may vary according to the problem at hand. The natural choice seems to be the connected components of the space of SAT-assignments, where two adjacent assignments have by definition Hamming distance 1. However, although this naive definition seems to work well on the satisfiability problem, it raises major difficulties on some other problems. For instance, in q -colorability, it is useful to permit color exchanges between two adjacent vertices in addition to single-vertex color changes. In XORSAT, the naive definition is inadequate, since jumps from solution to solution can involve a large, yet finite, Hamming distance due to the hard nature of linear Boolean constraints [36].

On the other hand, the existence of a gap in the x -satisfiability property is stronger than the original clustering hypothesis. Clusters are expected to have a typical size, and to be separated by a typical distance. However, even for typical formulas, there exist atypical clusters, the sizes and separations of which may differ from their typical values. Because of this variety of cluster sizes and separations, a large range of distances is available to pairs of SAT-assignments, which our x -satisfiability analysis takes into account. What we have shown suggests that, for typical formulas, the maximum size of all clusters is smaller than the minimum distance between two clusters (for a certain range of α and $K \geq 8$). This is a sufficient condition for clustering, but by no means a necessary one. As a matter of fact, our large K analysis conjectures that $\alpha_1(K)$ (the smaller α such that Conjecture 2 is verified) scales as $2^{K-1} \ln 2$, whereas $\alpha_d(K)$ (where the replica symmetry breaking occurs) and $\alpha_s(K)$ (where the one-step RSB Ansatz is supposed to be valid) scale as $2^K \ln K/K$ [21]. According to the physics interpretation, in the range $\alpha_s(K) < a < \alpha_1(K)$, there exist clusters, but they are not detected by the x -satisfiability approach. This limitation might account for the failure of our method for small values of K — even though more sophisticated techniques for evaluating the x -satisfiability threshold $\alpha_c(K, x)$ might yield some results for $K < 8$. Still, the conceptual simplicity of our method makes it a useful tool for proving similar phenomena in other systems of computational or physical interest.

A better understanding of the structure of the space of SAT-assignments could be gained by computing the average configurational entropy of pairs of clusters at fixed distance, which contains details about how intra-cluster sizes and inter-cluster distances are distributed. This would yield the value of the x -satisfiability threshold. Such a computation was carried out at a heuristic level within the framework of the cavity method for the random XORSAT problem [37], and should be extendable to the satisfiability problem or to other CSPs.

This work has been supported in part by the EC through the network MTR 2002-00319 ‘STIPCO’ and the FP6 IST consortium ‘EVERGROW’.

References

- [1] M. Sellitto, G. Biroli and C. Toninelli, *Facilitated spin models on Bethe lattice: Bootstrap percolation, mode-coupling transition and glassy dynamics*, Europhys. Lett. **69** (2005), 496–502.
- [2] J. Barré, A. R. Bishop, T. Lookman, A. Saxena, *On adaptability and “intermediate phase” in randomly connected networks*, Phys. Rev. Lett. **94**, 208701 (2005).
- [3] Robert G. Gallager. *Information Theory and Reliable Communication*, Wiley, New York, 1968.
- [4] David J.C. MacKay. *Information Theory, Inference & Learning Algorithms*, Cambridge University Press, Cambridge, 2002.
- [5] Stephen Cook. *The complexity of theorem proving procedures*, In Proceedings of the Third Annual ACM Symposium on Theory of Computing (1971), 151–158.
- [6] R. Monasson, R. Zecchina, *Statistical mechanics of the random K -satisfiability model*, Phys. Rev. E **56** (1997), 1357–1370.
- [7] T. Hogg, B. A. Huberman, C. P. Williams, *Phase transitions and the search problem*, Artificial Intelligence **81** (1996), 1–15.
- [8] Special Issue on *NP-hardness and Phase transitions*, edited by O. Dubois, R. Monasson, B. Selman and R. Zecchina, Theor. Comp. Sci. **265**, Issue: 1-2 (2001).
- [9] S. Kirkpatrick, B. Selman, *Critical Behavior in the Satisfiability of Random Boolean Expressions*, Science **264** (1994), 1297–1301.
- [10] R. Monasson, R. Zecchina, S. Kirkpatrick, B. Selman, and L. Troyanski, *Computational complexity from ‘characteristic’ phase transitions*, Nature **400** (1999), 133–137.
- [11] E. Friedgut, *Sharp Thresholds of Graph Properties, and the k -sat Problem*. J. Amer. Math. Soc. **12** (1999), no. 4, 1017–1054.
- [12] L. M. Kirousis, E. Kranakis, D. Krizanc, *A Better Upper Bound for the Unsatisfiability Threshold*, Technical report TR-96-09, School of Computer Science, Carleton University, 1996.

- [13] O. Dubois, Y. Boufkhad, *A general upper bound for the satisfiability threshold of random r -sat formulae*, J. Algorithms **24**(2) (1997), 395–420.
- [14] M.-T. Chao, J. Franco, *Probabilistic analysis of a generalization of the unit-clause literal selection heuristics for the k -satisfiability problem*, Inform. Sci. **51**(3) (1990), 289–314.
- [15] A. M. Frieze, S. Suen, *Analysis of two simple heuristics on a random instance of k -SAT*, J. Algorithms **20** (1996), 312–355.
- [16] D. Achlioptas, C. Moore, *The Asymptotic Order of the Random k -SAT Threshold*, Proc. Foundations of Computer Science (2002), 779–788.
- [17] D. Achlioptas, Y. Peres, *The Threshold for Random k -SAT is $2^k \log 2 - O(k)$* , Journal of the AMS, **17** (2004), 947–973.
- [18] M. Mézard, G. Parisi, *The Bethe lattice spin glass revisited*, Eur. Phys. J. **B 20** (2001), 217–233, and *The Cavity Method at Zero Temperature*, J. Stat. Phys. **111** (2003), 1–34.
- [19] M. Mézard, R. Zecchina, *Random K -satisfiability problem: From an analytic solution to an efficient algorithm*, Phys. Rev. E **66** (2002), 056126.
- [20] M. Mézard, G. Parisi, R. Zecchina, *Analytic and algorithmic solution of random satisfiability problems*, Science **297** (2002), 812–815.
- [21] S. Mertens, M. Mézard, R. Zecchina, *Threshold values of Random K -SAT from the cavity method*, Random Structures and Algorithms **28** (2006), 340–373.
- [22] A. Braunstein, M. Mezard, R. Zecchina, *Survey propagation: an algorithm for satisfiability*, Random Structures and Algorithms **27** (2005), 201–226.
- [23] A. Montanari, F. Ricci-Tersenghi, *On the nature of the low-temperature phase in discontinuous mean-field spin glasses*, Eur. Phys. J. **B 33** (2003), 339–346.
- [24] A. Montanari, G. Parisi, F. Ricci-Tersenghi, *Instability of one-step replica-symmetry-broken phase in satisfiability problems*, J. Phys. A **37** (2004), 2073–2091.
- [25] G. Semerjian, R. Monasson, *A Study of Pure Random Walk on Random Satisfiability Problems with “Physical” Methods*, Proceedings of the SAT 2003 conference, E. Giunchiglia and A. Tacchella eds., Lecture Notes in Computer Science (Springer) **2919** (2004), 120–134.
- [26] R. Mulet, A. Pagnani, M. Weigt, R. Zecchina, *Coloring Random Graphs*, Phys. Rev. Lett. **89** (2002), 268701.
- [27] A. Braunstein, R. Mulet, A. Pagnani, M. Weigt, R. Zecchina, *Polynomial iterative algorithms for coloring and analyzing random graphs*, Phys. Rev. E **68** (2003), 036702.
- [28] O. C. Martin, M. Mézard, O. Rivoire, *Frozen Glass Phase in the Multi-index Matching Problem*, Phys. Rev. Lett. **93** (2004), 217205.
- [29] A. Montanari, *The glassy phase of Gallager codes*, Eur. Phys. J. **B 23** (2001), 121–136.
- [30] S. Franz, M. Leone, A. Montanari, F. Ricci-Tersenghi, *Dynamic phase transition for decoding algorithms*, Phys. Rev. E **66** (2002), 046120.
- [31] M. Mézard, F. Ricci-Tersenghi, R. Zecchina, *Two Solutions to Diluted p -*

- Spin Models and XORSAT Problems*, J. Stat. Phys. **111** (2003), 505-533.
- [32] S. Cocco, O. Dubois, J. Mandler, R. Monasson, *Rigorous Decimation-Based Construction of Ground Pure States for Spin-Glass Models on Random Lattices*, Phys. Rev. Lett. **90** (2003), 047205.
- [33] O. Dubois, J. Mandler, *The 3-XORSAT threshold*, Proceedings of the 43th Annual IEEE Symposium on Foundations of Computer Science, Vancouver, pp. 769–778 (2002).
- [34] M. Mézard, T. Mora, R. Zecchina, *Clustering of solutions in the random satisfiability problem*, Phys. Rev. Lett. **94** (2005), 197205.
- [35] D. Achlioptas, F. Ricci-Tersenghi, *On the Solution-Space Geometry of Random Constraint Satisfaction Problems*, Proc. 38th annual ACM symposium on Theory of computing (2006), p. 130.
- [36] A. Montanari, G. Semerjian, *On the dynamics of the glass transition on Bethe lattices*, J. Stat. Phys. **124**, 103 (2006).
- [37] T. Mora, M. Mézard, *Geometrical organization of solutions to random linear Boolean equations*, J. Stat. Mech. (2006) P10007.

“Geometrical organization of solutions to
random linear Boolean equations”

J. Stat. Mech. (2006) P10007

Geometrical organization of solutions to random linear Boolean equations

Thierry Mora and Marc Mézard

Laboratoire de Physique Théorique et Modèles statistiques, UMR 8626, CNRS
and Université Paris Sud, Orsay Cedex, F-91405, France
E-mail: thierry.mora@u-psud.fr and marc.mezard@u-psud.fr

Received 5 September 2006

Accepted 26 September 2006

Published 16 October 2006

Online at stacks.iop.org/JSTAT/2006/P10007

doi:[10.1088/1742-5468/2006/10/P10007](https://doi.org/10.1088/1742-5468/2006/10/P10007)

Abstract. The random XORSAT problem deals with large random linear systems of Boolean variables. The difficulty of such problems is controlled by the ratio of number of equations to number of variables. It is known that in some range of values of this parameter, the space of solutions breaks into many disconnected clusters. Here we study precisely the corresponding geometrical organization. In particular, the distribution of distances between these clusters is computed by the cavity method. This allows one to study the ‘ x -satisfiability’ threshold, the critical density of equations where there exist two solutions at a given distance.

Keywords: cavity and replica method, message-passing algorithms, typical-case computational complexity

ArXiv ePrint: [cond-mat/0609099](https://arxiv.org/abs/cond-mat/0609099)

Contents

1. Introduction	2
2. Notation and definitions	4
3. Leaf removal as an instance of survey propagation	6
4. Distance landscape: thermodynamical approach	10
5. Diameter	13
6. Minimal and maximal distances between clusters	15
7. Conclusion and discussion	19
Acknowledgments	20
References	20

1. Introduction

Constraint satisfaction networks (CSN) are problems involving many discrete variables, with values in a finite alphabet, related by low density constraints: each constraint involves a finite number of variables. Such problems arise in many branches of science, from statistical physics (spin or structural glasses [1]) to information theory (low density parity check (LDPC) codes [2, 3]) and combinatorial optimization (satisfiability, colouring [4]). The ‘thermodynamic limit’ of such problems is obtained when the number of variables and the number of constraints go to infinity, keeping their ratio, the density of constraints α , fixed. A lot of attention has been focused in recent years on the study of random CSN, both because of their practical interest in coding, and also as a means to study ‘typical case’ complexity (as opposed to the traditional worst case complexity analysis). Many CSN are known to undergo a SAT–UNSAT phase transition when the density of constraints increases: there is a sharp threshold separating a SAT phase where all constraints can be satisfied with probability 1 in the thermodynamic limit from an UNSAT phase where, with probability 1, there is no configuration of the variables satisfying all the constraints. While the existence of a sharp threshold has been proved by Friedgut [5] for satisfiability and colouring, there is not yet any rigorous proof of the widely accepted conjecture according to which the threshold density of constraints converges to a fixed value α_c in the thermodynamic limit.

Recent years have seen an upsurge of statistical physics methods in the study of CSN. In particular, the replica method and the cavity method have been used to study the phase diagram [6]–[8]. Their most spectacular results are some arguably exact (but not yet rigorously proved) expressions for α_c and the existence of an intermediate SAT phase, in a region of constraint density $]\alpha_d, \alpha_c[$, where the space of solutions is split into many clusters, far away from each other. This clustering is an important building block of the theory: it is at the origin of the necessity to use the cavity method at the so-called

one-step replica symmetry breaking (1RSB) level; this method can be seen as a message-passing procedure and used as an algorithm for finding a SAT assignment of the variables. This algorithm, called survey propagation, turns out to be very powerful in satisfiability and colouring, and its effectiveness can be seen as one indirect piece of evidence in favour of clustering. On intuitive grounds, clustering is often held responsible for blocking many local search algorithms [9]. Although there does not exist any general discussion of this statement, this phenomenon was thoroughly investigated in the case of XORSAT [23].

The clustering effect can be studied in a more formal way by introducing the notion of x -satisfiability [10, 11]. A CSN with N variables is said x -satisfiable (x -SAT) if there exists a pair of SAT assignments of the variables which differ in a number of variables, $\in [Nx - \epsilon(N), Nx + \epsilon(N)]$. Here x is the reduced distance, which we keep fixed as N goes to infinity. The resolution $\epsilon(N)$ has to be sublinear in N : $\lim_{N \rightarrow \infty} \epsilon(N)/N = 0$, but its precise form is unimportant for our large N analysis. For example we can choose $\epsilon(N) = \sqrt{N}$. For many random CSN, it is reasonable to conjecture, in parallel with the existence of a satisfiability threshold, that x -satisfiability has a sharp threshold $\alpha_c(x)$ such that:

- if $\alpha < \alpha_c(x)$, a random formula is x -SAT almost surely;
- if $\alpha > \alpha_c(x)$, a random formula is x -UNSAT almost surely.

This conjecture has been proposed for k -satisfiability of random Boolean formulae where each clause involves exactly k variables with $k \geq 3$. So far only a weaker conjecture, analogous to Friedgut's theorem [5], has been established [11]. It states the existence of a non-uniform threshold $\alpha_c^{(N)}(x)$. Rigorous bounds on $\alpha_c(x)$ have been found in [11] for the k -satisfiability problem with $k \geq 8$, using moment methods developed in [12], but so far this x -satisfiability threshold has not been computed.

In this paper we compute the x -satisfiability threshold $\alpha_c(x)$ in the random XORSAT problem using the cavity method. This is a problem of random linear equations with Boolean algebra. It is important because many efficient error correcting codes are based on low density parity checks, the decoding of which involves precisely such linear systems. It is also one of the best understood cases of CSN. In particular, efforts to extend the replica method [13] and the cavity method [14] to deal with models defined on finite-connectivity lattices have resulted in the first exact (but non-rigorous) derivation of its phase diagram [15]. Later, a clear characterization of these clusters, combined with simple combinatoric arguments, gave a rigorous basis to these predictions [16]–[18]. These works have computed the phase diagram in detail and provide expressions for the two thresholds $\alpha_d < \alpha_c < 1$.

Our computation of $\alpha_c(x)$ confirms this known structure, and it also provides insight into the geometrical structure of clusters. We find that $\alpha_c(x)$ is non-monotonic (see figure 5), which confirms the existence of gaps in distances where there do not exist any pairs of solutions.

The method used in our computation is in itself interesting. It turns out that it is not possible to compute $\alpha_c(x)$ directly, by fixing x and varying α . Instead, we work at a fixed value of α and introduce a probability distribution for pairs of SAT assignments, where the distance between the solutions plays the role of the energy. The computation of the entropy as a function of the energy, and more precisely the computation of the energies where it vanishes, then allows one to reconstruct $\alpha_c(x)$. Our computation thus involves a

mixture of hard constraints (the fact that the two assignments must satisfy the XORSAT formula) and soft constraints (the Boltzmann weight which depends on their distance). This is reflected in the structure of the cavity fields that solve this problem.

The remainder of this paper is organized as follows. The next section introduces some notation. In section 3, we analyse classical survey propagation on XORSAT and show its equivalence with the ‘leaf removal’ [18] or ‘decimation’ [16] algorithm. This analysis allows one to re-derive the phase diagram of XORSAT and sets up useful notation and concepts for later computations. In section 4 we perform a statistical mechanics analysis of weight properties in a single cluster using the cavity method. Section 5 applies this formalism to the computation of the cluster diameter, while section 6 is devoted to the evaluation of inter-cluster distances. In section 7 we sum up and discuss our results.

2. Notation and definitions

An XORSAT formula is defined on a string of N variables $x_1, x_2, \dots, x_N \in \{0, 1\}$ by a set of M parity checks of the form

$$\sum_{i \in V(a)} x_i = y_a \pmod{2}, \quad \text{for all } a = 1, \dots, M \quad (1)$$

where $y_a \in \{0, 1\}$. Here $V(a) \subset \{1, \dots, N\}$ is the subset of variables involved in parity check a . Later on $i \in a$ will be used as shorthand for $i \in V(a)$.

Equation (1) can be rewritten in the matrix form

$$A\mathbf{x} = \mathbf{y} \pmod{2}, \quad A = \{A_{ia}\}_{i \in [N], a \in [M]} \quad (2)$$

where $A_{ia} = 1$ if $i \in a$ and $A_{ia} = 0$ otherwise. The pair $F = (A, \mathbf{y})$ defines the formula. Such a linear system can be solved in polynomial time by Gaussian elimination. If a formula has solutions, it is SAT; otherwise, it is UNSAT. The thermodynamics limit is $N \rightarrow \infty$, $M \rightarrow \infty$ with a fixed density of constraints $\alpha = M/N$.

In this paper we specialize to random k -XORSAT formulae, where each equation involves a subset of k variables, chosen independently with uniform probability among the $\binom{N}{k}$ possible ones, and each y_a independently takes value 0 or 1 with probability $1/2$. One important characterization of a XORSAT formula $F = (A, \mathbf{y})$ is the number $\mathcal{N}_N(F)$ of assignments of the Boolean variables $\mathbf{x}$ which satisfy all the equations, and the corresponding entropy density

$$s_N(F) = \frac{1}{N} \log \mathcal{N}_N(F) \quad (3)$$

The logarithm is base 2 throughout the paper. Using a spin representation $\sigma_i = (-1)^{x_i}$, the k -XORSAT problem can also be mapped onto a spin glass model where interactions involve products of k spins (the variables $(-1)^{y_a}$ then play the role of quenched random exchange couplings) [15], and the question of whether a formula is SAT is equivalent to asking whether the corresponding spin glass instance is frustrated.

Previous work [15]–[18] has shown that:

- For $\alpha < \alpha_d(k)$, the formula is SAT, almost surely (i.e. with probability $\rightarrow 1$ as $N \rightarrow \infty$). The solution set forms one big connected component, and the entropy density concentrates at large N to $(N - M)/N = 1 - \alpha$; this phase is called the EASY-SAT phase.

- For $\alpha_d(k) < \alpha < \alpha_c(k)$, the formula is still SAT almost surely, but the solution set is made of an exponentially large (in N) number of components far away from each other (in the following we shall give a precise definition of these clusters); the entropy density also concentrates at large N to $(N - M)/N = 1 - \alpha$. This is the HARD-SAT phase.
- For $\alpha > \alpha_c(k)$ (with $\alpha_c(k) < 1$), the formula is UNSAT almost surely. The entropy is $-\infty$. This second transition is the usual SAT-UNSAT transition.

The fact that, throughout the SAT phase ($\alpha < \alpha_c(k)$), the entropy density concentrates to $1 - \alpha$ is not surprising: it can be understood as the fact that matrix A has rank M almost surely in the SAT phase. The intuitive reason is that, each time there exists a linearly dependent set of checks, the choice of y_a has probability $1/2$ of leading to a contradiction. So the rank of A cannot differ much from M in the SAT phase. From the point of view of linear algebra, the existence of the clustered phase, i.e. the fact that the vector subspace of SAT assignments breaks into disconnected pieces, is more surprising, as is the discontinuity of $s_N(F)$ at the transition α_c . These two aspects are in fact related: the quantity which vanishes at the SAT-UNSAT transition is actually the log of the number of clusters of solutions, while each cluster keeps a finite volume.

We will study the geometric properties of the space of solutions for random k -XORSAT in the HARD-SAT phase using the notion of x -satisfiability. In terms of solutions of linear equations, we want to know whether there exist two Boolean vectors $\mathbf{x}$ and $\mathbf{x}'$ which both satisfy $A\mathbf{x} = A\mathbf{x}' = \mathbf{y}$, where the Hamming distance $d_{\mathbf{x},\mathbf{x}'} \equiv (\mathbf{x} - \mathbf{x}')^2 = Nx$. Clearly, if such a pair exists, $\mathbf{x} - \mathbf{x}'$ is a solution to the homogeneous ('ferromagnetic') problem where $\mathbf{y} = \mathbf{0}$:

$$A(\mathbf{x} - \mathbf{x}') = \mathbf{0}. \quad (4)$$

Therefore, a formula $F = (A, \mathbf{y})$ is x -SAT if and only if F is SAT and if there exists a solution $\mathbf{x}$ to the homogeneous system $A\mathbf{x} = \mathbf{0}$ of weight $d_{\mathbf{x},\mathbf{0}} \approx Nx$ (the *weight* is by definition the distance to $\mathbf{0}$). Note that for $x = 0$, this second condition is automatically fulfilled and x -satisfiability is equivalent to satisfiability. This linear space structure also implies that the set of solutions looks the same seen from any solution in the SAT phase: the number of solutions at distance d of any given solution $\mathbf{x}_0$ is independent from $\mathbf{x}_0$.

Distance properties can also be investigated directly by evaluating extremal distances between solutions. To that end we define three distances: (a) the cluster diameter d_1 , i.e. the largest Hamming distance between solutions belonging to the same cluster; this diameter is independent of the cluster; (b) the minimal and maximal inter-cluster distances d_2 and d_3 , i.e. the smallest and largest, respectively, Hamming distance between solutions belonging to distinct clusters. All three distances are assumed to be self-averaging in the thermodynamic limit of the random problem: $x_1(\alpha) = d_1/N$, $x_2(\alpha) = d_2/N$ and $x_3(\alpha) = d_3/N$ will denote the corresponding limits. In the particular case where k is even, the formula is invariant under the transformation $\mathbf{x} \leftrightarrow \mathbf{x} + \mathbf{1} \pmod{2}$, which is reflected in terms of distances by a symmetry with respect to $x = 1/2$: $x \leftrightarrow 1 - x$. A direct consequence is that $x_3(\alpha) = 1 - x_2(\alpha)$, and that a fourth weight, defined as $1 - x_1(\alpha)$, will also come into play. These distance functions are related to the x -satisfiability threshold as follows: at fixed α , a formula is x -SAT almost surely iff

- $x \in [0, x_1(\alpha)] \cup [x_2(\alpha), x_3(\alpha)]$ when k is odd;

- $x \in [0, x_1(\alpha)] \cup [x_2(\alpha), 1 - x_2(\alpha)] \cup [1 - x_1(\alpha), 1]$ when k is even.

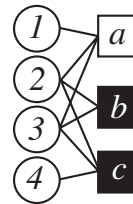
We will now compute x_1, x_2, x_3 with the cavity method.

3. Leaf removal as an instance of survey propagation

XORSAT formulae are conveniently represented by factor graphs, called *Tanner graphs*, in which variables and checks form two distinct types of node, with the simple rule that the edge (i, a) between i and a is present if $i \in a$.

An example of a Tanner graph and its associated linear system is shown below:

- (a) $x_1 + x_2 + x_3 = 0 \pmod{2}$
 (b) $x_2 + x_3 = 1 \pmod{2}$
 (c) $x_2 + x_3 + x_4 = 1 \pmod{2}$.



The number of variables involved in a check a , denoted by $|V(a)|$, is the degree of a in the factor graph. Here we study k -XORSAT where this degree is fixed at k . Similarly, if $V(i)$ denotes the set of parity checks in which i is represented, $|V(i)|$ is the degree of i in the factor graph. The degrees of checks are commonly referred to as *right-degrees* and those of variables as *left-degrees*. The infinite-length (thermodynamic) limit is obtained by sending N and M to infinity while keeping the ratio $\alpha = M/N$ fixed. In this limit, the distribution of left-degrees is a Poisson law of parameter $k\alpha$: the probability of a variable having degree ℓ is $\pi_{k\alpha}(\ell)$, where $\pi_x(\ell) = \exp(-x)x^\ell/\ell!$.

Here we use the *leaf removal algorithm* (LR) in order to obtain a precise definition of the notion of ‘cluster’ or ‘component’ of solutions, one which is valid also for finite N . The algorithm proceeds as follows: pick a variable of degree 1 (called a *leaf*), remove it as well as the only check it is connected to. Continue the process until there remains no leaf. The interest of this algorithm is easily seen: a variable on a leaf can always be assigned in such a way that the (unique) check to which it is connected is satisfied.

The linear system remaining after leaf removal is independent of the order in which leaves are removed. It is called the *core*. A ‘core check’ is a check which only involves core variables. If the core is empty, the problem is trivially SAT. In general, given a solution of the core, one can easily reconstruct a solution of the complete formula by running leaf removal in the reverse direction, in a scheme which we refer to as *leaf reconstruction*. In this procedure, checks are added one by one along with their leaves, starting from the core. If an added check involves only one leaf, the value of that variable is determined uniquely so that the check is satisfied. If the number of leaves k' is greater than 1, one can choose the joint value of those leaves among $2^{k'-1}$ possibilities. The process is iterated until the complete factor graph has been rebuilt. Given a core solution, one can construct many solutions to the complete formula. Variables which are uniquely determined by the core solution are called *frozen*, and variables that can fluctuate are called *floppy*. Of course, by definition, the frozen part includes the core itself. A core solution defines a *cluster*. All solutions built from the same core solution belong to the same cluster. We shall see later how this definition fits in the intuitive picture that we sketched previously in terms of connectedness.

We propose here an alternative to the leaf removal algorithm, which also builds the core, but keeps actually more information. The approach is inspired by the cavity method, and is a special instance of survey propagation (SP) [7]. To each edge (i, a) one assigns two numbers $\hat{m}_{a \rightarrow i}^t$ and $m_{i \rightarrow a}^t$ belonging to $\{0, 1\}$, updated as follows:

- At $t = 0$, $\hat{m}_{a \rightarrow i}^0 = 1$, $m_{i \rightarrow a}^0 = 1$ for all edges (i, a) .
- $m_{i \rightarrow a}^{t+1} = 1 - \prod_{b \in i-a} (1 - \hat{m}_{b \rightarrow i}^t)$.
- $\hat{m}_{a \rightarrow i}^t = \prod_{j \in a-i} m_{j \rightarrow a}^t$.
- Stop when $\hat{m}_{a \rightarrow i}^{t+1} = \hat{m}_{a \rightarrow i}^t$ for all (i, a) .

Here $a \in i$ is a shorthand for $a \in V(i)$.

The interpretation of $m_{i \rightarrow a}^t = 1$ is: ‘variable i is constrained at time t in the absence of check a ’, and $\hat{m}_{a \rightarrow i}^t = 1$: ‘check a constrains variable i at time t ’. One also defines $M_i^t = 1 - \prod_{a \in i} (1 - \hat{m}_{a \rightarrow i}^t) \in \{0, 1\}$. This number indicates whether node i is constrained at time t ($M_i^t = 1$) or not ($M_i^t = 0$).

At $t = 0$, all variables are constrained. The algorithm consists in detecting the underconstrained variables and propagating the information through the graph to simplify the formula. At the first step, only variables of degree 1 are affected: if i is of degree 1 and is connected to a , $m_{i \rightarrow a}^1 = 1 - \prod_{\emptyset} = 0$. This, in turn, gives freedom to a , which no longer constrains its other variables: $\hat{m}_{a \rightarrow j}^1 = 0$, for $j \in a - i$. This effectively removes a and i from the formula, just as in the leaf removal algorithm. In the subsequent steps of the iteration, there will be considered as a *leaf* (in the LR sense) a variable i such that there exists exactly one $a \in i$ such that $\hat{m}_{a \rightarrow i}^t = 1$. In that case we have $m_{i \rightarrow a}^{t+1} = 0$, thus implementing a step of LR.

Let us add a word about the term ‘survey propagation’ we have used so far. Analysis of the 1RSB cavity equations at zero temperature [18] (see [7] for a more complete discussion in the case of k -SAT) shows that cavity biases fall into two categories, depending on the edge we consider: either a warning is sent (compelling taking the value 0 or 1 depending on the cluster, with probability a half for each), or no warning is sent. (In more technical terms, the survey propagation reduces to warning propagation.) The first situation corresponds in our language to $\hat{m}_{a \rightarrow i} = 1$ and the second to $\hat{m}_{a \rightarrow i} = 0$. Similarly, we have $m_{i \rightarrow a} = 1$ if the cavity field is non-zero and $m_{i \rightarrow a} = 0$ otherwise. Therefore our algorithm carries the same information as survey propagation.

The interest of SP over leaf removal is that it keeps track of the leaves which are uniquely determined by their check. For example, if two or more leaves are connected to the same check a at time t , at time $t + 1$ one has $\hat{m}_{a \rightarrow i}^{t+1} = 0$ for all $i \in a$, reflecting the fact that a cannot uniquely determine the value of several leaves. Conversely, if a is connected to a unique leaf i and if one has $m_{j \rightarrow a}^t = 1$ for all $j \in a - i$, then one gets $\hat{m}_{a \rightarrow i}^t = 1$, reflecting the fact that, the variables $\{x_j\}_{j \in a-i}$ being fixed in the absence of a , i is determined uniquely.

A little reasoning shows that when the algorithm stops ($t = t_f$), i is frozen iff $M_i^{t_f} = 1$, and i belongs to the core iff there exist at least two checks $a, b \in i$ such that $\hat{m}_{a \rightarrow i}^{t_f} = \hat{m}_{b \rightarrow i}^{t_f} = 1$. In the final state, we say that the directed edge $i \rightarrow a$ is frozen if $m_{i \rightarrow a} \equiv m_{i \rightarrow a}^{t_f} = 1$ and that $a \rightarrow i$ is frozen if $\hat{m}_{a \rightarrow i} \equiv \hat{m}_{a \rightarrow i}^{t_f} = 1$. In the opposite case, edges are called floppy (see figure 1). This version of SP is strictly equivalent to the *belief propagation* algorithm used for decoding low density parity check codes on the binary erasure channel, also called the ‘peeling decoder’ in that context.

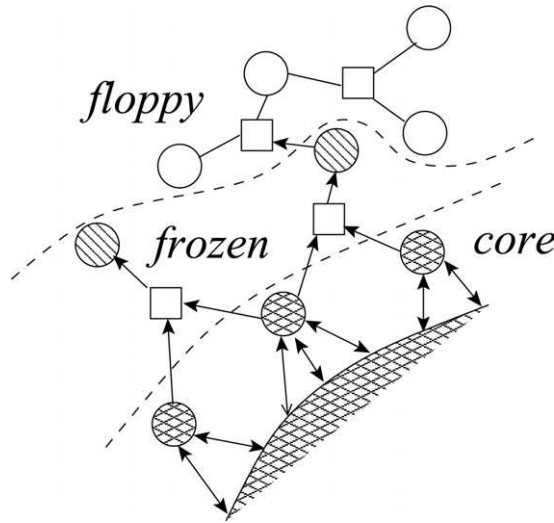


Figure 1. An example of a fixed point of SP. Circles represent variable nodes, and squares check nodes. An arrow means that message m or $\hat{m}$ has value 1, that is, that the directed edge is frozen when SP stops. Leaf removal propagates null messages from the outer leaves down to the core, while ‘leaf reconstruction’ propagates non-null messages from the core up the frozen part.

SP can be studied by density evolution in order to derive the phase diagram, as in [18]. Let us briefly survey this study for completeness. The statistics of messages at time t is described by two numbers:

$$v^t = \frac{1}{Mk} \sum_{(i,a)} \delta(m_{i \rightarrow a}^t, 0), \quad w^t = \frac{1}{Mk} \sum_{(i,a)} \delta(\hat{m}_{a \rightarrow i}^t, 0), \quad (5)$$

where the sums run over all edges of the Tanner graph. When $N \rightarrow \infty$, these densities are governed by evolution equations:

$$\begin{aligned} v^{t+1} &= \sum_{\ell} \pi_{k\alpha}(\ell) (w^t)^\ell = \exp[-k\alpha(1 - w^t)] \\ w^t &= 1 - [1 - v^t]^{k-1}, \end{aligned} \quad (6)$$

which are initialized with $v^0 = w^0 = 0$. These equations are exact if the Tanner graph is a tree. In our case the graph is locally tree-like (it is a tree up to finite distance when seen from a generic point) and one could set up a rigorous proof of (6) using the methods developed in [19].

The fixed point of these equations is given by the *cavity equation*:

$$w = 1 - \{1 - e^{-k\alpha(1-w)}\}^{k-1}. \quad (7)$$

Setting $\lambda = k\alpha(1 - w)$, equation (7) can be rewritten as

$$\lambda = k\alpha(1 - e^{-\lambda})^{k-1} \quad (8)$$

When $\alpha < \alpha_d$, the unique fixed point is $\lambda = 0$ (i.e. $w = 1$). This means that the core is empty. For $\alpha > \alpha_d$ however, there remains an extensive core of size

$$N_c = N \left[\sum_{\ell \geq 2} \pi_{k\alpha}(\ell) (1 - w^\ell - \ell w^{\ell-1}) \right] = N [1 - (1 + \lambda)e^{-\lambda}] \quad (9)$$

while the number of frozen variables is

$$N_f = N \left[\sum_{\ell \geq 2} \pi_{k\alpha}(\ell) (1 - w^\ell) \right] = N [1 - e^{-\lambda}]. \quad (10)$$

The number of core checks is

$$M_c = M(1 - v)^k = \alpha N [1 - e^{-\lambda}]^k. \quad (11)$$

The left-degree distribution (with respect to core checks) inside the core is given by a truncated Poissonian:

$$P_c(\ell) = \frac{1}{e^\lambda - 1 - \lambda} \frac{\lambda^\ell}{\ell!} \mathbb{I}(\ell \geq 2), \quad (12)$$

where $\mathbb{I}$ is the indicator function.

One can show that the leaf removal algorithm conserves the uniformity of the ensemble. Therefore, the core formula is a random XORSAT formula with right-degree k and left-degree distribution $P_c(\ell)$ given by (12). The number of solutions to such a formula is known to concentrate to its mean value when the size goes to infinity [17, 18]. In the case of the core formula, this number is simply $2^{N_c - M_c}$ if $N_c \geq M_c$ and 0 otherwise. Recalling that the complete formula has solutions if and only if the core formula does, we find that the SAT-UNSAT threshold α_c is given by the equation

$$1 - (1 + \lambda)e^{-\lambda} = \alpha [1 - e^{-\lambda}]^k. \quad (13)$$

The number of clusters is characterized by the *complexity* or *configurational entropy*, that is the logarithm of the number of core solutions:

$$\Sigma(\alpha) = \frac{1}{N} \log(\# \text{clusters}) = \frac{N_c - M_c}{N} = 1 - (1 + \lambda)e^{-\lambda} - \alpha [1 - e^{-\lambda}]^k. \quad (14)$$

We recall that the group structure of the solution set implies that all clusters have the same internal structure. Their common internal entropy is therefore given by

$$s_{\text{inter}} = 1 - \alpha - \Sigma(\alpha) \quad (15)$$

where we have used the fact that the total entropy is $1 - \alpha$.

Let us comment on the relationship between our definition of clusters and the more traditional one. Usually, clusters are defined as the ‘connected’ components of the solution set, where connectedness is to be understood in the following way: two solutions are connected if one can go from one to the other by a sequence of solutions separated by a finite Hamming distance (when $N \rightarrow \infty$). To make contact with our own definition of clusters, one needs to prove two things. First, that two solutions built from the same core solution are connected. Second, that two core solutions are necessarily separated by an extensive Hamming distance ($\geq cN$, with c constant), which implies that two solutions built from two distinct core solutions are not connected. Both proofs can be found in [18]. This reconciles our definition (which holds for any single instance of XORSAT) with the usual one (which only makes sense for infinite-length ensembles).

4. Distance landscape: thermodynamical approach

As we have already observed, studying pairs of solutions is equivalent to studying solutions to the ferromagnetic problem. Indeed, if S denotes the affine subspace of solutions to $A\mathbf{x} = \mathbf{y}$, and S_0 the vector subspace of solutions to $A\mathbf{x} = \mathbf{0}$, we have

$$S \times S = \{(\mathbf{x}', \mathbf{x}' + \mathbf{x}), (\mathbf{x}', \mathbf{x}) \in S \times S_0\}. \quad (16)$$

In particular, distances in S are reflected by weights in S_0 . Therefore, in order to study the range of attainable distances between solutions, one just needs to study the range of possible weights in S_0 . To that end we set a thermodynamical framework in which the weight plays the role of an energy:

$$E(\mathbf{x}) \equiv |\mathbf{x}| = \sum_i \delta_{x_i, 1}. \quad (17)$$

The Boltzmann measure at temperature β^{-1} is thus defined by

$$\mathbb{P}(\mathbf{x}, \beta) = \frac{1}{Z(\beta)} \prod_a \delta_{\mathbb{F}_2} \left(\sum_{i \in a} x_i, 0 \right) 2^{-\beta|\mathbf{x}|} \quad (18)$$

where the normalization constant $Z(\beta)$ is the partition function. The Dirac delta function, here defined on the two-element field $\mathbb{F}_2$, enforces that only configurations of S_0 are considered. Remarkably, this measure is formally similar to the one used to infer the most probable codeword under maximum-likelihood decoding in low density parity check (LDPC) codes on the binary symmetric channel [20]. In fact, as we shall see soon, some of the methods used to solve both problems share common aspects.

A very useful scheme for estimating marginal probabilities in models defined on sparse graphs is the cavity method [14], which we have already mentioned in the previous section. Let $p_{i \rightarrow a}^x$ be the probability that $x_i = x$ under the measure defined by (18), where the link (i, a) has been removed. The replica symmetric (RS) cavity method consists in computing the cavity marginals $p_{i \rightarrow a}^x$ (viewed as variable-to-check messages) using a closed set of equations where check-to-variable messages are also introduced as intermediate quantities. These second-kind messages are denoted by $q_{a \rightarrow i}^x$ and are proportional to the probability that $x_i = x$ when i is connected to a only. Messages are updated until convergence occurs with the following rules:

$$p_{i \rightarrow a}^{x_i} = \frac{1}{Z_{i \rightarrow a}} \prod_{b \in i-a} q_{b \rightarrow i}^{x_i} 2^{-\beta \delta_{x_i, 1}} \quad (19)$$

$$q_{a \rightarrow i}^{x_i} = \sum_{\{x_j\}_{j \in a-i}} \prod_{j \in a-i} p_{j \rightarrow a}^{x_j} \delta_{\mathbb{F}_2} \left(\sum_{j \in a} x_j, 0 \right) \quad (20)$$

where $Z_{i \rightarrow a}$ is a normalization constant. When convergence is reached, marginal probabilities are obtained as

$$p_i^{x_i} \equiv \sum_{\{x_j\}_{j \neq i}} \mathbb{P}(\mathbf{x}, \beta) = \frac{1}{Z_{i+a \in i}} \prod_{a \in i} q_{a \rightarrow i}^{x_i} 2^{-\beta \delta_{x_i, 1}} \quad (21)$$

where $Z_{i+a \in i}$ is also a normalization constant. Continuing the analogy with codes, it is interesting to note that these cavity equations are identical [21] to the belief propagation

(BP) equations [22] used to decode messages with LDPC codes on the binary symmetric channel.

It turns out that cavity equations (19), (20) do not admit a unique solution, as one would expect if the system were replica symmetric. Instead, let us show that they admit exactly one solution for each cluster. In a given cluster denoted by $\mathbf{c}$, let us denote by c_i the value of a frozen variable i . There exists a solution to (19), (20), where, for every frozen variable i ,

$$\begin{aligned} p_{i \rightarrow a}^x &= \delta_{x, c_i} & \text{if } i \rightarrow a \text{ frozen,} \\ q_{a \rightarrow i}^x &= \delta_{x, c_i} & \text{if } a \rightarrow i \text{ frozen.} \end{aligned} \quad (22)$$

In order to show that this is a solution, let us use the SP messages, which provide information on how the fixing of the core solution forces the values of frozen variables. For example $m_{i \rightarrow a} = 1$ indicates that x_i is entirely determined by the core solution, supposing that the edge (i, a) has been removed. Consider the SP fixed point relations

$$\begin{aligned} \hat{m}_{a \rightarrow i} &= \prod_{j \in a-i} m_{j \rightarrow a}, \\ m_{i \rightarrow a} &= 1 - \prod_{b \in i-a} (1 - \hat{m}_{b \rightarrow i}). \end{aligned} \quad (23)$$

They are in fact *contained* in the cavity equations (19), (20). In fact, the iteration of cavity equations allows one to identify the frozen edges, irrespectively of the cluster the system falls into.

But the cavity equations also contain ‘fluctuating’ messages, where p^x and q^x are in $]0, 1[$, which are *de facto* restricted to the floppy part. We parametrize them by the cavity fields and biases:

$$\beta h_{i \rightarrow a}^c = \log \frac{p_{i \rightarrow a}^0}{p_{i \rightarrow a}^1}, \quad \beta u_{a \rightarrow i}^c = \log \frac{q_{a \rightarrow i}^0}{q_{a \rightarrow i}^1} \quad (24)$$

which satisfy the equations

$$h_{i \rightarrow a}^c = \sum_{b \in i-a} u_{b \rightarrow i}^c + 1 \quad \text{with } i \rightarrow a \text{ floppy,} \quad (25)$$

$$\beta u_{a \rightarrow i}^c = 2 \operatorname{arctanh} \left[\prod_{j \in a^{nf}-i} \tanh(\beta h_{j \rightarrow a}^c / 2) \prod_{j \in a^f-i} (-1)^{c_j} \right] \quad \text{with } a \rightarrow i \text{ floppy,} \quad (26)$$

where a^f (resp. a^{nf}) is the set of neighbours i of a such that $i \rightarrow a$ is frozen (resp. floppy). Note that cavity messages $h_{i \rightarrow a}^c$ and $u_{a \rightarrow i}^c$ now depend explicitly on the cluster considered, and are uniquely determined by it.

The multiplicity of solutions to RS cavity equations is a clear sign that the replica symmetry is broken. The main lesson from this discussion is that solutions can fluctuate according to two hierarchical levels of statistics: the first level deals with fluctuations inside a single cluster, i.e. fluctuations on the floppy part, while the second level deals with the choice of the cluster. The reduced cavity equations (25), (26) correctly describe the first level¹, when the system is forced to live in cluster $\mathbf{c}$. This leads to defining a

¹ Although the RS ansatz is unable to describe the whole system, it can reasonably be assumed to be valid on a single cluster.

new probability measure and partition function, restricted to $\mathbf{c}$:

$$Z_{\mathbf{c}}(\beta) = \sum_{\mathbf{x} \in \mathbf{c}} 2^{-\beta \sum_{i=1}^N \delta_{x_i, 1}}. \quad (27)$$

By construction, this system is characterized by the fixing of the frozen edges (22) and by the reduced cavity equations (25), (26). The second level of statistics, i.e. the statistics over the clusters, is appropriately handled by a 1RSB calculation and will be the subject of section 6. We first focus on the properties of single clusters under the measure defined by (27).

The cavity method comes with a technique for estimating the log of the partition functions, also called the potential in our case:

$$\phi(\beta) = -\frac{1}{N} \log Z(\beta). \quad (28)$$

(Note that this quantity differs from the usual free energy by a factor β .) It can be computed within the RS ansatz using the Bethe formula [21]:

$$N\phi(\beta) = \sum_i \Delta\phi_{i+a \in i} - (k-1) \sum_a \Delta\phi_a \quad (29)$$

where

$$\begin{aligned} \Delta\phi_{i+a \in i} &= -\log Z_{i+a \in i} = -\log \sum_{x_i} \prod_{a \in i} q_{a \rightarrow i}^{x_i} 2^{-\beta \delta_{x_i, 1}} \\ \Delta\phi_a &= -\log \sum_{\{x_i\}_{i \in a}} \prod_{i \in a} p_{i \rightarrow a}^{x_i} \delta_{\mathbb{F}_2} \left(\sum_{i \in a} x_i, 0 \right). \end{aligned} \quad (30)$$

This formula has a rather simple interpretation: $\Delta\phi_{i+a \in i}$ is the contribution of i and its adjacent checks to the potential. When these contributions are summed, each check is counted k times, whence the need to subtract $k-1$ times the contribution of each check $\Delta\phi_a$. Also note that this expression is variational: it is stationary in the messages $\{p_{i \rightarrow a}\}$ as soon as the cavity equations (19), (20) are satisfied.

The RS ansatz is valid in a single cluster. The single cluster potential $\phi_{\mathbf{c}}(\beta) = -(1/N) \log Z_{\mathbf{c}}(\beta)$ can therefore be computed by plugging equations (22), (25) and (26) into the Bethe formula (30), provided one uses the messages corresponding to one given cluster $\mathbf{c}$. When one is restricted to a single cluster $\mathbf{c}$, the range of possible weights is $[x_{\mathbf{c}}, X_{\mathbf{c}}]$. The minimal and maximal weights can be obtained by sending $\beta \rightarrow \pm\infty$. For $\beta \rightarrow \infty$, the second cavity equation (26) simplifies to

$$u_{a \rightarrow i}^{\mathbf{c}} = \mathcal{S} \left(\prod_{j \in a^{nf} - i} h_{j \rightarrow a}^{\mathbf{c}} \prod_{j \in a^f - i} (-1)^{c_j} \right) \min_{j \in a^{nf} - i} |h_{j \rightarrow a}^{\mathbf{c}}| \quad \text{with } a \rightarrow i \text{ floppy} \quad (31)$$

where $\mathcal{S}(x) = 1$ if $x > 0$, -1 if $x < 0$ and 0 if $x = 0$.

The ‘ground state energy’, i.e. the minimal weight in $\mathbf{c}$, is obtained as

$$x_{\mathbf{c}} = \lim_{\beta \rightarrow \infty} \partial_{\beta} \phi_{\mathbf{c}}(\beta) = \frac{1}{N} \sum_{i \text{ floppy}}^N \frac{1 - \mathcal{S}(\sum_{a \in i} u_{a \rightarrow i}^{\mathbf{c}} + 1)}{2} + \frac{1}{N} \sum_{i \text{ frozen}} \delta_{c_i, 1}. \quad (32)$$

The $\beta \rightarrow -\infty$ limit yields very similar equations. These equations will be analysed in the next section.

Let us also write down the equations giving the potential, which will be used in section 6:

$$N\phi_{\mathbf{c}}(\beta) = \sum_i \Delta\phi_{i+a \in i}^{\mathbf{c}} - (k-1) \sum_a \Delta\phi_a^{\mathbf{c}} \quad (33)$$

$$\lim_{\beta \rightarrow \infty} \frac{1}{\beta} \Delta\phi_{i+a \in i}^{\mathbf{c}} \equiv \Delta x_{i+a \in i}^{\mathbf{c}}, \quad \lim_{\beta \rightarrow \infty} \frac{1}{\beta} \Delta\phi_a^{\mathbf{c}} \equiv \Delta x_a^{\mathbf{c}} \quad \text{with} \quad (34)$$

$$\Delta x_{i+a \in i}^{\mathbf{c}} = \frac{1}{2} \left(\sum_{a \in i} |u_{a \rightarrow i}^{\mathbf{c}}| + 1 - \left| \sum_{a \in i} u_{a \rightarrow i}^{\mathbf{c}} + 1 \right| \right) \quad \text{if } i \text{ is floppy} \quad (35)$$

$$\Delta x_{i+a \in i}^{\mathbf{c}} = \sum_{a \in i^{nf}} |u_{a \rightarrow i}^{\mathbf{c}}| \vartheta(-u_{a \rightarrow i}^{\mathbf{c}}) \quad \text{if } i \text{ is frozen and } c_i = 0 \quad (36)$$

$$\Delta x_{i+a \in i}^{\mathbf{c}} = 1 + \sum_{a \in i^{nf}} |u_{a \rightarrow i}^{\mathbf{c}}| \vartheta(u_{a \rightarrow i}^{\mathbf{c}}) \quad \text{if } i \text{ is frozen and } c_i = 1 \quad (37)$$

$$\Delta x_a^{\mathbf{c}} = \vartheta \left(- \prod_{i \in a^{nf}} h_{i \rightarrow a}^{\mathbf{c}} \prod_{i \in a^f} (-1)^{c_i} \right) \min_{i \in a^{nf}} |h_{i \rightarrow a}^{\mathbf{c}}|, \quad (38)$$

where i^f and i^{nf} are defined in a similar fashion to a^f and a^{nf} .

5. Diameter

With our formalism, computing the cluster diameter boils down to computing the maximal weight in cluster $\mathbf{0}$ (the cluster containing $\mathbf{0}$). The relevant partition function for this task is

$$Z_{\mathbf{0}}(\beta) = 2^{-N\phi_{\mathbf{0}}(\beta)} = \sum_{\mathbf{x} \in \mathbf{0}} \delta_{\mathbb{F}_2} \left(\sum_{i \in a} x_i, 0 \right) 2^{-\beta \sum_{i=1}^N \delta_{x_i, 1}}. \quad (39)$$

When $\beta \rightarrow -\infty$, the solution of the cavity equations corresponding to cluster $\mathbf{0}$ is characterized by

$$\begin{aligned} p_{i \rightarrow a}^x &= \delta_{x, 0} && \text{if } i \rightarrow a \text{ frozen,} \\ q_{a \rightarrow i}^x &= \delta_{x, 0} && \text{if } a \rightarrow i \text{ frozen,} \\ h_{i \rightarrow a} &= \sum_{b \in i-a} u_{b \rightarrow i} + 1 && \text{if } i \rightarrow a \text{ floppy,} \\ u_{a \rightarrow i} &= -\mathcal{S} \left[\prod_{j \in a^{nf}-i} (-h_{j \rightarrow a}) \right] \min_{j \in a^{nf}-i} |h_{j \rightarrow a}| && \text{if } a \rightarrow i \text{ floppy} \end{aligned} \quad (40)$$

and the maximum weight d_1 is given by

$$d_1 = \lim_{\beta \rightarrow -\infty} \partial_{\beta} \phi_{\mathbf{0}}(\beta) = \sum_{i \text{ floppy}}^N \frac{1 + \mathcal{S}(\sum_{a \in i} u_{a \rightarrow i} + 1)}{2}. \quad (41)$$

These equations are presented for single XORSAT formulae, and can be solved by simple iteration of the corresponding message-passing rules. In practice however, in the

Geometrical organization of solutions to random linear Boolean equations

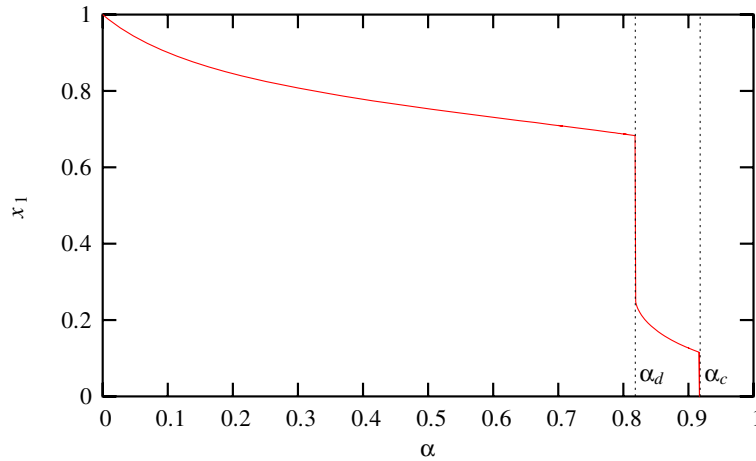


Figure 2. Diameter of a cluster of solutions. When one decreases α below α_d all clusters aggregate into one big cluster, thus explaining the discontinuity.

regime where α is near (but smaller than) α_d , one does not always reach convergence. This is arguably due to the hard nature of XORSAT constraints, as was pointed out in [23]: as one nears the dynamical transition, hopping from one solution to the other requires an increasing (yet sub-extensive) number of changes, making the sampling of solutions difficult. To circumvent this problem, we can work directly in the infinite-length limit by considering the probability distribution functions (pdfs) of each kind of message:

$$\begin{aligned} P(h) &= \frac{1}{Mk} \sum_{(i,a)} \delta_{h, h_{i \rightarrow a}} \\ Q(u) &= \frac{1}{Mk} \sum_{(i,a)} \delta_{u, u_{a \rightarrow i}}. \end{aligned} \quad (42)$$

When $N \rightarrow \infty$, self-consistency equations for these distributions read

$$\begin{aligned} P(h) &= \sum_{\ell} \pi_{k\alpha w}(\ell) \int \prod_{a=1}^{\ell} du_a Q(u_a) \delta \left(h - \sum_{a=1}^{\ell} u_a - 1 \right) \\ Q(u) &= \frac{1}{w} \sum_{i=1}^{k-1} \binom{k-1}{i} v^i (1-v)^{k-1-i} \int \prod_{j=1}^i dh_j P(h_j) \delta \left[u + \mathcal{S} \left(\prod_{j=1}^i (-h_j) \right) \min_j |h_j| \right] \end{aligned} \quad (43)$$

and one has

$$x_1(\alpha) = \lim_{N \rightarrow \infty} \frac{d_1}{N} = e^{-\lambda} \int dh P(h) \frac{1 + \mathcal{S}(h)}{2}. \quad (44)$$

These equations can be solved with a population dynamics algorithm [14]. In figure 2, we represent the maximal diameter x_1 as a function of α .

6. Minimal and maximal distances between clusters

In section 4 we have set up the formalism for computing the minimal and the maximal weights in a given cluster $\mathbf{c}$ using the cavity method. In order to evaluate the minimal and maximal weights in *all* clusters expect $\mathbf{0}$, we resort to a statistical treatment of the cavity equations. This scheme is known as the 1RSB cavity method in the replica language. We first specialize to the case of minimal weights, the other case being formally equivalent. We already know that the number of clusters grows exponentially with N . Here we further assume that the number of clusters with a given minimal weight $x_{\mathbf{c}}$ is exponential in N , and we define the complexity

$$\sum_{\mathbf{c} \neq \mathbf{0}} \delta(x, x_{\mathbf{c}}) = 2^{N\Sigma_m(x)}. \quad (45)$$

With this quantity we associate the 1RSB potential

$$2^{N\psi_m(y)} = \sum_{\mathbf{c} \neq \mathbf{0}} 2^{-Nyx_{\mathbf{c}}} = \int dx \, 2^{N(\Sigma_m(x) - yx)}. \quad (46)$$

When N is large, a saddle-point evaluation of this quantity yields

$$\psi_m(y) = \min_x [yx - \Sigma_m(x)] = yx^* - \Sigma_m(x^*) \quad \text{with } y = \partial_x \Sigma_m(x^*) \quad (47)$$

and $\psi_m(y)$ is thus related to $\Sigma_m(x)$ by a Legendre transformation. In terms of statistical mechanics, m is an inverse temperature coupled to the ‘energy’ $x_{\mathbf{c}}$; the complexity plays the role of a microcanonical entropy, and the potential is equivalent to a free energy, up to a factor m . The minimal weight in all clusters (expect $\mathbf{0}$) is given by the smallest x such that $\Sigma_m(x) \geq 0$. Our goal is now to compute $\psi_m(y)$ and to infer $\Sigma_m(x)$ by inverse Legendre transformation.

We proceed to the statistical analysis of the cavity equations under Boltzmann measure $2^{-Nyx_{\mathbf{c}}}$. This amounts to writing 1RSB cavity equations, where messages are distributions of RS messages over all clusters. The distribution of messages on floppy edges is described by the two pdfs:

$$P^{i \rightarrow a}(h) = \langle \delta(h, h_{i \rightarrow a}^{\mathbf{c}}) \rangle \quad (48)$$

$$Q^{a \rightarrow i}(u) = \langle \delta(u, u_{a \rightarrow i}^{\mathbf{c}}) \rangle. \quad (49)$$

The average $\langle \cdot \rangle$ is performed with the aforementioned measure on clusters, with the implicit assumption that the edge (i, a) has been removed. On frozen edges, messages are trivial, but their values depend on the cluster considered. We thus define for frozen edges

$$P_0^{i \rightarrow a} = \langle \delta(p_{i \rightarrow a}^0, 1) \rangle \quad P_1^{i \rightarrow a} = 1 - P_0^{i \rightarrow a} \quad (50)$$

$$Q_0^{a \rightarrow i} = \langle \delta(q_{a \rightarrow i}^0, 1) \rangle \quad Q_1^{a \rightarrow i} = 1 - Q_0^{a \rightarrow i}. \quad (51)$$

In order to write a closed set of equations for these probability distributions, we need to know how the Boltzmann weight $2^{-Nyx_{\mathbf{c}}}$ biases the message-passing procedure: when a field $h_{i \rightarrow a}$ is estimated as a function of its ‘grandparents’ ($\{h_{j \rightarrow b}\}$, $j \in b - i$, $b \in i - a$), a reweighting term $2^{-y\Delta x_{i \rightarrow a}}$ is associated with it [7, 14], where $\Delta x_{i \rightarrow a}$ is the contribution

of i and its adjacent checks (except a) to the total weight. This contribution is obtained as $\Delta x_{i+a \in i}$ in equations (35)–(37), but with a removed.

The 1RSB cavity equations read

- $i \rightarrow a$ frozen:

$$P_0^{i \rightarrow a} = \frac{1}{\mathcal{Z}_{i \rightarrow a}} \prod_{b \in i^f - a} Q_0^{b \rightarrow i} \int \prod_{b \in i^{nf} - a} du_{b \rightarrow i} Q^{b \rightarrow i}(u_{b \rightarrow i}) 2^{-y \sum_{b \in i^{nf} - a} |u_{b \rightarrow i}|^{\vartheta} (-u_{b \rightarrow i})}$$

$$P_1^{i \rightarrow a} = \frac{1}{\mathcal{Z}_{i \rightarrow a}} \prod_{b \in i^f - a} Q_1^{b \rightarrow i} \int \prod_{b \in i^{nf} - a} du_{b \rightarrow i} Q^{b \rightarrow i}(u_{b \rightarrow i}) 2^{-y(1 + \sum_{b \in i^{nf} - a} |u_{b \rightarrow i}|^{\vartheta} (u_{b \rightarrow i}))},$$

- $i \rightarrow a$ floppy:

$$P^{i \rightarrow a}(h) = \frac{1}{\mathcal{Z}_{i \rightarrow a}} \int \prod_{b \in i - a} du_{b \rightarrow i} Q^{b \rightarrow i}(u_{b \rightarrow i}) 2^{-y/2(\sum_{b \in i - a} |u_{b \rightarrow i}| + 1 - |\sum_{b \in i - a} u_{b \rightarrow i}| + 1)}$$

$$\times \delta \left(h - 1 - \sum_{b \in i - a} u_{b \rightarrow i} \right)$$

(here and in the previous equations $\mathcal{Z}_{i \rightarrow a}$ is a normalization constant),

- $a \rightarrow i$ frozen:

$$Q_0^{a \rightarrow i} = \frac{1 + \prod_{j \in a - i} (2P_0^{j \rightarrow a} - 1)}{2},$$

- $a \rightarrow i$ floppy:

$$Q^{a \rightarrow i}(u) = \sum_{\substack{\{c_j=0,1\} \\ j \in a^f - i}} \prod_{j \in a^f - i} P_{c_j}^{j \rightarrow a} \int \prod_{j \in a^{nf} - i} dh_{j \rightarrow a} P^{j \rightarrow a}(h_{j \rightarrow a})$$

$$\times \delta \left[u - \mathcal{S} \left(\prod_{j \in a^{nf} - i} h_{j \rightarrow a} \prod_{j \in a^f - i} (-1)^{c_j} \right) \min_{j \in a^{nf} - i} |h_{j \rightarrow a}| \right].$$

The potential $\psi_m(y)$ is obtained by a Bethe-like formula [7]:

$$N\psi_m(y) = \sum_i \Delta\psi_{i+a \in i} - (k-1) \sum_a \Delta\psi_a \quad (56)$$

with

$$\begin{aligned} \Delta\psi_{i+a \in i} &= -\log \langle 2^{-y\Delta x_{i+a \in i}} \rangle = -\log \mathcal{Z}_{i+a \in i} \\ \Delta\psi_a &= -\log \langle 2^{-y\Delta x_a} \rangle \\ &= -\log \frac{1 + \prod_{i \in a} (2P_0^{i \rightarrow a} - 1)}{2} \quad \text{if } a \in \text{core} \\ &= -\log \sum_{\substack{\{c_i=0,1\} \\ i \in a^f}} \prod_{j \in a^f} P_{c_i}^{i \rightarrow a} \int \prod_{i \in a^{nf}} dh_{i \rightarrow a} P^{i \rightarrow a}(h_{i \rightarrow a}) \\ &\quad \times \exp \left[-y \log(2) \vartheta \left(- \prod_{i \in a^{nf}} h_{i \rightarrow a} \prod_{i \in a^f} (-1)^{c_i} \right) \min_{i \in a^{nf}} |h_{i \rightarrow a}| \right] \quad \text{otherwise} \end{aligned} \quad (57)$$

where $\mathcal{Z}_{i+a \in i}$ is defined as $\mathcal{Z}_{i \rightarrow a}$ but in the presence of a .

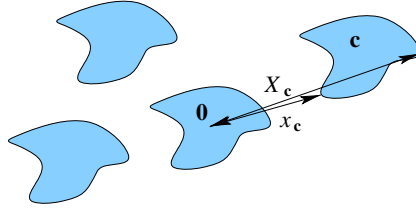


Figure 3. Pictorial representation of the clustered space of solutions around $\mathbf{0}$ in the N -dimensional hypercube. For a cluster $\mathbf{c}$, the minimal and maximal distances $x_{\mathbf{c}}$ and $X_{\mathbf{c}}$ are depicted.

Like in the diameter calculation, 1RSB cavity equations can be interpreted as message-passing update rules, with the difference that messages are now surveys over all clusters. The output of that procedure is the minimal distance complexity $\Sigma_m(x)$, obtained as the inverse Legendre transform of $\psi_m(y)$. We refer to the corresponding algorithm as ‘distance survey propagation’. The same procedure can be implemented in the $\beta \rightarrow -\infty$ limit and yields the maximal distance complexity:

$$\Sigma_M(x) = \frac{1}{N} \log \sum_{\mathbf{c} \neq \mathbf{0}} \delta(x, X_{\mathbf{c}}), \quad (58)$$

where $X_{\mathbf{c}}$ is the maximal weight in cluster $\mathbf{c}$ (see figure 3). Note that in the particular case where $y = 0$, which corresponds to a uniform measure over the clusters, classical SP is recovered for both versions of the algorithm (minimal and maximal distance): in that limit we have $Q_0^{a \rightarrow i} = P_0^{i \rightarrow a} = 1/2$ and the calculation of $\psi_m(0)$ and $\psi_M(0)$ gives back $-\Sigma(\alpha)$, the total complexity (14), as expected.

The practical implementation of distance-SP demands particular care when small distances are considered: it turns out that distance complexities $\Sigma_m(x)$ and $\Sigma_M(x)$ are not concave, which entails that the functions $\psi_m(y)$ and $\psi_M(y)$ are multivalued in a certain range of y . A way to circumvent this problem (already used in [24]) is to keep the weight $x = \partial_y \psi_m(y)$ fixed after each iteration and to deduce y accordingly. Here is how the algorithm proceeds for a given reduced weight x :

- (1) Run classical SP.
- (2) Initialize all floppy and frozen messages $\{P_{i \rightarrow a}\}$, $\{Q_{a \rightarrow i}\}$ to random values. Choose a (reasonable) value for y .
- (3) Until convergence is reached, do:
 - Update all $a \rightarrow i$ messages $\{Q_{a \rightarrow i}\}$ and then all $i \rightarrow a$ messages $\{P_{i \rightarrow a}\}$ at inverse temperature y .
 - Find y such that $x = \partial_y \psi_m(y, \{P_{i \rightarrow a}\}, \{Q_{a \rightarrow i}\})$ by the secant method, $\{P_{i \rightarrow a}\}$ and $\{Q_{a \rightarrow i}\}$ being fixed.
- (4) Compute $\psi_m(y, \{P_{i \rightarrow a}\}, \{Q_{a \rightarrow i}\})$ as well as its derivative and deduce $\Sigma_m(x) = yx - \psi_m(y)$.

Note that since the messages are pdfs themselves, the update of each of them in step 3 is performed by a population dynamics subroutine.

Geometrical organization of solutions to random linear Boolean equations

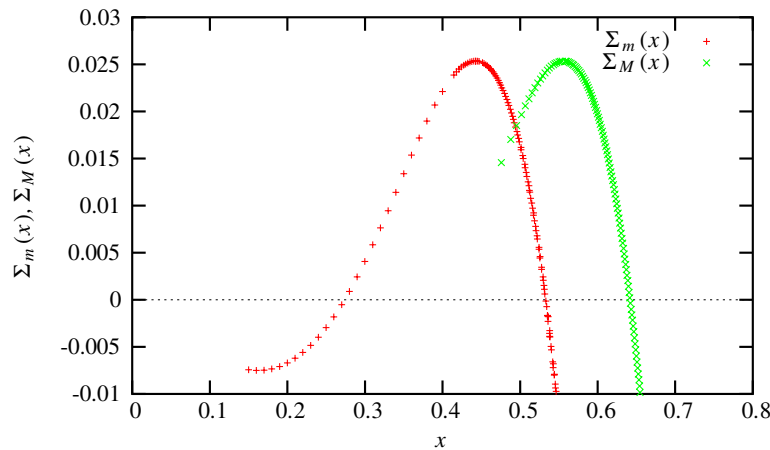


Figure 4. Minimal and maximal distance complexities as a function of the reduced distance x , for $k = 3$, $N = 10\,000$ and $M = 8600$.

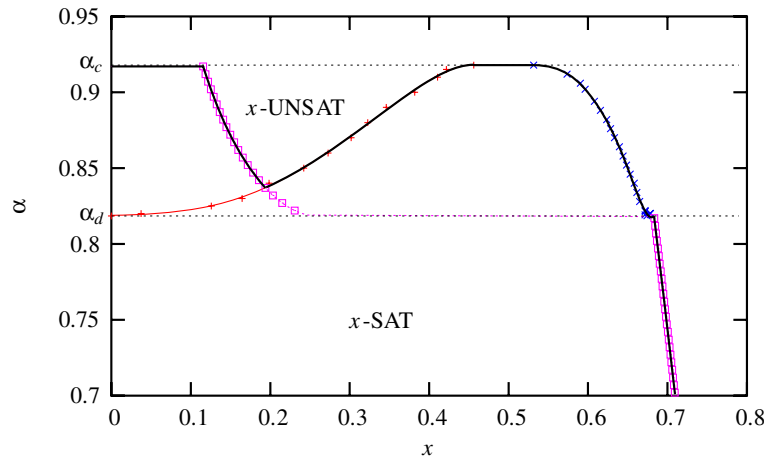


Figure 5. Phase diagram of the 3-XORSAT problem in the (x, α) plane. The cluster diameter ($\square$), as well as minimal (+) and maximal ($\times$) distances between solutions of distinct clusters, are represented. The thick line is the x -satisfiability threshold.

Figure 4 shows the minimal and maximal weight complexities $\Sigma_m(x)$ and $\Sigma_M(x)$ for a random 3-XORSAT formula with $N = 10\,000$ and $M = 8600$. These complexities can be regarded as kinds of weight enumerator functions for clusters. Their fluctuations from formula to formula can be significant (15%), even for large system sizes ($N = 10\,000$).

An average version (density evolution) of distance-SP can also be implemented for random k -XORSAT, in the same spirit as equation (43). Such a computation involves distributions (on edges) of distributions (on clusters) and can be solved by population dynamics, where each element of the population is itself a population. The zeros of $\overline{\Sigma_m(x)}$ and $\overline{\Sigma_M(x)}$ thus obtained yield the minimal and maximal inter-cluster distances $x_2(\alpha)$ and $x_3(\alpha)$, respectively, as shown in figure 5. Together with the cluster diameter $x_1(\alpha)$ computed in section 5, these values are used to construct the x -satisfiability threshold.

Our algorithm can in principle be run on any system of Boolean linear equations and is expected to give reasonable results provided that the loops of the underlying Tanner graph are large. The case of LDPC codes is of particular interest because it allows several simplifications and has been extensively studied from both the combinatorial [25] and statistical physics [24, 26] point of view. LDPC codes are homogeneous Boolean linear systems where parity checks and variables may have arbitrary degree distributions, with the restriction that variables should always have degrees no less than 2. This implies that the leaf removal algorithm is inefficient on such linear systems: all variables belong to the core, and are frozen. In particular, each cluster is made of one unique solution: the cluster diameter is 0, and the minimal and maximal inter-cluster distances coincide. Their common complexity $\Sigma_m(x) = \Sigma_M(x)$ is often called the ‘weight enumerator exponent’ and is an important property of ensembles of codes. Translated into our formalism, this means that all messages are frozen and the distance-SP algorithm simplifies dramatically:

$$P_0^{i \rightarrow a} = \frac{1}{\mathcal{Z}_{i \rightarrow a}} \prod_{b \in i^f - a} Q_0^{b \rightarrow i}, \quad P_1^{i \rightarrow a} = \frac{1}{\mathcal{Z}_{i \rightarrow a}} \prod_{b \in i^f - a} Q_1^{b \rightarrow i} 2^{-y} \quad (59)$$

$$Q_0^{a \rightarrow i} = \frac{1 + \prod_{j \in a - i} (2P_0^{j \rightarrow a} - 1)}{2}. \quad (60)$$

Not surprisingly, the density evolution analysis of this simplified algorithm yields the same equations as those obtained with the replica method in [24, 26].

7. Conclusion and discussion

We have applied the cavity method to estimate extremal distances between solutions of random linear systems with large girth in the clustered phase. Our results are used to compute the x -satisfiability threshold of the random k -XORSAT problem. The notion of x -satisfiability, which tells us whether one can find a pair of solutions separated by a Hamming distance x , was introduced in the context of another constraint satisfaction problem, k -SAT, where it was used to give rigorous evidence in favour of the clustering phenomenon [10].

Although k -XORSAT is a rather simple problem, it displays a very similar phase diagram to harder problems such as k -SAT and q -colourability. In particular, its clustered phase is well defined and understood. That said, finding extremal distances in the solution space of linear Boolean equations is a hard task in general: for instance, the decision problem associated with finding the minimal weight of LDPC codes is NP-complete [27].

We were able to compute three quantities: the cluster diameter, as well as the minimal and maximal inter-cluster distances. We believe our method to give a good approximation for systems with large girth and to be exact in the thermodynamic limit for random XORSAT. In the line of survey propagation, we devised a series of algorithms for these tasks, which explicitly exploit the clustered structure of the solution space. More precisely, the space of solutions is characterized by two hierarchical levels of fluctuations: inside and between clusters. In k -XORSAT, these two kinds of fluctuations are carried by two disjoint sets of variables, and our algorithms explicitly distinguish between these two kinds of variables. In the special case of LDPC codes, the point-like nature of clusters much

simplifies the equations, and previous expressions for the weight enumerator exponent obtained by the replica method are recovered.

The method presented here offers a number of generalizations. In particular, it could be used at finite temperature to yield the full weight enumerator function. More interestingly, it could be adapted to deal with other CSN, such as k -SAT, for which only bounds are known; unfortunately, numerical computations are in that case much heavier, albeit formally similar. Let us mention that a similar approach was followed in [28] in the case of q -colourability, with the difference that distances were estimated from a reference configuration (which is not a solution) instead of considering distances between solutions.

Our work studies the geometrical properties of the solution space by taking explicitly into account fluctuations inside clusters, captured by the ‘evanescent fields’. This very general approach, already explored in [28], allows one to gain a better understanding of the fine structure of the clustered phase and seems to us a promising direction for future work. Also, with similar tools, decimation schemes such as the one introduced in [7] could be used to select solutions or clusters with particular properties.

Acknowledgments

We would like to thank Andrea Montanari for sharing the numerical trick used in the replica evaluation of the weight enumerator function of LDPC codes [24]. This work has been supported in part by the EU through the network MTR 2002-00319 ‘STIPCO’ and the FP6 IST consortium ‘EVERGROW’.

References

- [1] Mézard M, Parisi G and Virasoro M A, 1987 *Spin-glass Theory and Beyond (Lecture Notes in Physics vol 9)* (Singapore: World Scientific)
- [2] Gallager R G, *Low-density parity check codes*, 1962 *IRE Trans. Inf. Theory* **8** 21
- [3] MacKay D J C, 2003 *Information Theory, Inference, and Learning Algorithms* (Cambridge: Cambridge University Press)
- [4] Papadimitriou C H, 1994 *Computational Complexity* (Reading, MA: Addison-Wesley)
- [5] Friedgut E, *Sharp thresholds of graph properties, and the k -SAT problem*, 1999 *J. Am. Math. Soc.* **12** 1017
- [6] Mézard M, Parisi G and Zecchina R, *Analytic and algorithmic solution of random satisfiability problems*, 2002 *Science* **297** 812–5
- [7] Mézard M and Zecchina R, *Random k -satisfiability problem: from an analytic solution to an efficient algorithm*, 2002 *Phys. Rev. E* **66** 056126
- [8] Mulet R, Pagnani A, Weigt M and Zecchina R, *Coloring random graphs*, 2002 *Phys. Rev. Lett.* **89** 268701
- [9] Semerjian G and Monasson R, *A study of pure random walk on random satisfiability problems with ‘physical’ methods*, 2004 *Proc. SAT 2003 Conf. (Lecture Notes in Computer Science vol 120)* ed E Giunchiglia and A Tachella (Berlin: Springer) p 2919
- [10] Mézard M, Mora T and Zecchina R, *Clustering of solutions in the random satisfiability problem*, 2005 *Phys. Rev. Lett.* **94** 197205
- [11] Mora T, Mézard M and Zecchina R, *Pairs of SAT assignments and clustering in random Boolean formulae*, 2005 *Preprint cond-mat/0506053*
- [12] Achlioptas D and Peres Y, *The threshold for random k -SAT is $2^k \log 2 - O(k)$* , 2004 *J. Am. Math. Soc.* **17** 947–73
- [13] Monasson R, *Optimization problems and replica symmetry breaking in finite connectivity spin-glasses*, 1998 *J. Phys. A: Math. Gen.* **31** 515
- [14] Mézard M and Parisi G, *The Bethe lattice spin glass revisited*, 2001 *Eur. Phys. J. B* **20** 217
- [15] Ricci-Tersenghi F, Weigt M and Zecchina R, *Simplest random k -satisfiability problem*, 2001 *Phys. Rev. E* **63** 026702
- [16] Cocco S, Dubois O, Mandler J and Monasson R, *Rigorous decimation-based construction of ground pure states for spin glass models on random lattices*, 2003 *Phys. Rev. Lett.* **90** 047205

- [17] Dubois O and Mandler J, *The 3-XORSAT threshold*, 2002 *Proc. 43rd Ann. IEEE Symp. on Foundations of Computer Science (FOCS '02)* p 769
- [18] Mézard M, Ricci-Tersenghi F and Zecchina R, *Alternative solutions to diluted p-spin models and XORSAT problems*, 2003 *J. Stat. Phys.* **111** 505
- [19] Richardson T and Urbanke R, *Modern Coding Theory*, 2006 at press, available at lthcwww.epfl.ch/mct
- [20] Nishimori H, 2001 *Statistical Physics of Spin Glasses and Information Processing: An Introduction* (Oxford: Oxford University Press)
- [21] Yedidia J S, Freeman W F and Weiss Y, *Constructing free energy approximations and generalized belief propagation algorithms*, 2002 Technical Report TR-2002-35, Mitsubishi Electrical Research Laboratories available at <http://www.merl.com>
- [22] Kschischang F R, Frey B and Loeliger H-A, *Factor graphs and the sum-product algorithm*, 2001 *IEEE Trans. Inf. Theory* **47** 498–519
- [23] Montanari A and Semerjian G, *On the dynamics of the glass transition on Bethe lattices*, 2005 Preprint cond-mat/0509366
- [24] Di C, Montanari A and Urbanke R, *Weight distributions of LDPC code ensembles: combinatorics meets statistical physics*, 2004 *Int. Symp. on Information Theory* (Piscataway, NJ: IEEE)
- [25] Di C, Proietti D, Telatar I E, Urbanke R L and Richardson T J, *Finite length analysis of low-density parity-check codes on the binary erasure channel*, 2002 *IEEE Trans. Inf. Theory* **48** 1570–9
- [26] Condamin S, *Study of the weight enumerator function for a Gallager code*, 2002 <http://www.inference.phy.cam.ac.uk/condamin/report.ps>
- [27] Vardy A, *The intractability of computing the minimum distance of a code*, 1997 *IEEE Trans. Inf. Theory* **43** 1757–66
- [28] Mézard M, Palassini M and Rivoire O, *Landscape of solutions in constraint satisfaction problems*, 2005 *Phys. Rev. Lett.* **95** 200202

“Error Exponents of Low-Density Parity-Check
Codes on the Binary Erasure Channel”

IEEE Information Theory Workshop,
2006 (ITW '06), Chengdu. pp. 81–85

Error Exponents of Low-Density Parity-Check Codes on the Binary Erasure Channel

Thierry Mora

Laboratoire de Physique Théorique et
Modèles Statistiques, Bât. 100
Université Paris-Sud and CNRS
F-91405 Orsay, France.
Email: mora@lptms.u-psud.fr

Olivier Rivoire

Laboratory of Living Matter
The Rockefeller University
1230 York Avenue, Box 34
New York, NY-10021, USA
Email: orivoire@rockefeller.edu

Abstract — We introduce a thermodynamic (large deviation) formalism for computing error exponents in error-correcting codes. Within this framework, we apply the heuristic cavity method from statistical mechanics to derive the average and typical error exponents of low-density parity-check (LDPC) codes on the binary erasure channel (BEC) under maximum-likelihood decoding.

I. INTRODUCTION

Assessing the performance of error-correcting codes is a founding topics of information theory. Amongst the simplest codes are the binary *block codes*, where a source generates with equal probability one of 2^L *codewords*, each a sequence of N bits. As a codeword is transmitted through a *discrete memoryless channel*, a noise ξ alters independently each bit with some probability. The *binary erasure channel* (BEC), for instance, erases a bit with a prescribed probability $p \in [0, 1]$. Given the received message, the decoding task consists in inferring the most likely original codeword. The probability of error $\mathbb{P}_\xi(\text{error}|\mathcal{C}_N)$ then provides a simple characterization of the performance of a code $\mathcal{C}_N$.

The properties of error-correcting codes are conveniently studied through *ensembles* of codes $\mathcal{C}_N$, consisting for instance of the set of all block codes with length N and *rate* $R = L/N$. Shannon showed that, in the limit $N \rightarrow \infty$, a typical code in such an ensemble has a vanishing probability of error if (and only if) $R < R_c(p)$, where $R_c(p)$ corresponds to the *channel capacity*. This capacity is simply $R_c(p) = 1 - p$ for the BEC. We are here interested in refining the description of the error probability beyond the channel capacity. *Error exponents* give the exponential rate of decay of $\mathbb{P}_\xi(\text{error}|\mathcal{C}_N)$ with N , for $\mathcal{C}_N \in \mathcal{C}_N$, and offer the most appealing generalization. Of particular interest is the so-called *reliability function*, which gives the lowest achievable exponents as a function of the rate R [2]. However, despite significant efforts to estimate error exponents, resulting in the establishment of a number of bounds, exact expressions are scarce and restricted to a few extreme cases.

In this note, we put forward a *thermodynamic* (or *large deviation*) formalism [13] for evaluating error exponents in error-correcting codes. This formalism coherently encompasses two types of exponents: if $\mathcal{C} = \{\mathcal{C}_N\}_{N \geq 1}$ de-

notes a sequence of ensembles of codes, we can indeed define, depending on the procedure for choosing the codes $\mathcal{C}_N$ in the ensembles $\mathcal{C}_N$, an *average* and a *typical* error exponents as

$$E_{\text{av}} = - \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E}_{\mathcal{C}_N} [\mathbb{P}_\xi(\text{error}|\mathcal{C}_N)], \quad (1)$$

$$E_{\text{typ}} = - \lim_{N \rightarrow \infty} \frac{1}{N} \mathbb{E}_{\mathcal{C}_N} [\log \mathbb{P}_\xi(\text{error}|\mathcal{C}_N)], \quad (2)$$

where $\mathbb{E}_{\mathcal{C}_N}$ denotes the expectation value when $\mathcal{C}_N$ is drawn uniformly from the ensemble $\mathcal{C}_N$ (log is base 2 throughout). Although the typical error exponent is the most interesting from the practical point of view, the average error exponent is usually simpler to estimate theoretically.

We analyze in the thermodynamic formalism one of the most promising family of block codes, the low-density parity-check (LDPC) codes [5]. The codewords of these codes correspond to the kernel of a sparse $M \times N$ *parity-check matrix* A , with $M = N - L$. Different choices for A lead to different ensemble of codes $\mathcal{C}_N$, the simplest example being regular ensembles¹ defined with A having ℓ 1's per column and k per line, and zeros otherwise (in which case $R = 1 - \ell/k$). LDPC codes have been shown to formally map to physical models of disordered systems on random graphs [7], and we shall exploit this analogy to apply the (non-rigorous) cavity method [12] recently proposed in this context² (see also [14] for a related approach).

II. THERMODYNAMIC FORMALISM

Given a received word, consisting of a codeword from a code $\mathcal{C}_N$ altered by a noise ξ on the BEC, let $\mathcal{N}_N(\xi, \mathcal{C}_N)$ be the number of codewords from which it could come from (this quantity is independent of the initial codeword with LDPC codes). By definition, decoding is achievable if and only if $\mathcal{N}_N(\xi, \mathcal{C}_N) = 1$. For random codes, the geometry of the space of codewords indicates that, at least in

¹In this paper we restrict to regular codes, even though our method can be generalized to any irregular ensemble [11].

²While the exponential scaling of the error probability is guaranteed when the ensemble of codes comprises all block codes, the average error probability of LDPC codes is known to be polynomial in N [5]. Following Gallager, we shall ignore the few atypical codes responsible for this behavior, and consider the average error exponent associated with an expurgated ensemble where they have been excluded [5].

the vicinity of the channel capacity, an error most probably involves an exponential number of potential codewords (see e.g. [1]). In such situations, we characterize $\mathcal{N}_N(\xi, \mathcal{C}_N)$ by an *entropy*, defined as

$$S_N(\xi, \mathcal{C}_N) = \log \mathcal{N}_N(\xi, \mathcal{C}_N). \quad (3)$$

In the limit $N \rightarrow \infty$, for sequences of codes $\mathcal{C} = \{\mathcal{C}_N\}_N$ taken from the sequence of ensembles $\mathcal{C} = \{\mathcal{C}_N\}_N$, the entropy density $s = S_N/N$ concentrates to a well defined value $\bar{s}$, and the channel coding theorem takes the following form: there exists p_c , such that $\bar{s} = 0$ for $p < p_c$, and $\bar{s} > 0$ for $p > p_c$ [4]. More generally, we postulate that, for a typical sequence of codes $\mathcal{C}^0 = \{\mathcal{C}_N^0\}_N$, the entropy S_N satisfies a *large deviation principle* [3], i.e.,

$$\mathbb{P}_\xi[S_N(\xi, \mathcal{C}_N^0)/N = s] \asymp 2^{-NL_0(s)}, \quad (4)$$

with $a_N \asymp b_N$ meaning that $\log a_N / \log b_N \rightarrow 1$. The typical value $\bar{s}$ corresponds here to the minimum of the *rate function* L_0 , with $L_0(\bar{s}) = 0$. In cases where L_0 is strictly convex, the typical error exponent is obtained as

$$\begin{aligned} E_{\text{typ}} &= -\lim_{N \rightarrow \infty} \frac{1}{N} \log \sum_{s \geq 1/N} \mathbb{P}_\xi[S_N(\xi, \mathcal{C}_N^0)/N = s] \\ &= L_0(s = 0). \end{aligned} \quad (5)$$

A simpler quantity to compute than $L_0(s)$ is $L_1(s)$, the rate function for the large deviations of $S_N(\xi, \mathcal{C}_N)$ with respect to both the noise ξ and the codes $\mathcal{C}_N$,

$$\mathbb{P}_{\xi, \mathcal{C}_N}[S_N(\xi, \mathcal{C}_N)/N = s] \asymp 2^{-NL_1(s)}. \quad (6)$$

In the so-called *thermodynamic formalism* [13], $L_1(s)$ is associated with a *potential* $\phi(x)$ defined through the relation

$$2^{N\phi(x)} = \mathbb{E}_{\xi, \mathcal{C}_N}[2^{xS_N(\xi, \mathcal{C}_N)}] \asymp \int ds \, 2^{N[xs - L_1(s)]}. \quad (7)$$

Under the assumption that it is convex, the rate function $L_1(s)$ is derived from the knowledge of $\phi(x)$ by Legendre transformation:

$$L_1(s) = \max_x [xs - \phi(x)]. \quad (8)$$

The average exponent, obtained from $E_{\text{av}} = L_1(s = 0)$, may differ from the typical exponent E_{typ} . Typical codes $\mathcal{C}_N^0$ can however also be described within a thermodynamic formalism, provided an extra “temperature” y is introduced, together with a generalized potential $\psi(x, y)$ satisfying

$$2^{N\psi(x, y)} = \mathbb{E}_{\mathcal{C}_N} \left[\left(\mathbb{E}_\xi[2^{xS_N(\xi, \mathcal{C}_N)}] \right)^y \right]. \quad (9)$$

The average case is here recovered for $y = 1$, with $\psi(x, y = 1) = \phi(x)$. Typical error exponents are associated with $y = 0$ (see [11] for details and exceptions), with

$$E_{\text{typ}} = L_0(s = 0) = -\partial_y \psi(x^*, y = 0), \quad (10)$$

where x^* selects for $s = \frac{1}{y} \partial_x \psi(x^*, y) \Big|_{y=0} = 0$.

III. CAVITY METHOD

Disordered systems constructed out of random ensembles, of which LDPC codes are particular examples, have been the subject of intensive studies in statistical mechanics. One of the most elaborate analytical tool developed in this context is the *cavity method* [10], which allows to extract the typical properties of models defined on random graphs. While yielding virtually equivalent predictions than the similar *replica method*, this method has both more sound probabilistic foundations, and an attractive relation to message-passing algorithms, such as *belief propagation* (BP). The cavity method has also been recently extended to deal with large deviations [12], making it perfectly suited to the evaluation of error exponents.

As far as typical codes and typical noise are concerned, the cavity method is equivalent to a BP *density evolution* analysis. Belief propagation, also known as the “peeling decoder” in the context of the BEC [8], consists in propagating messages between *bits* (the N letters of a word) and *checks* (the M linear equations encoded in the parity-check matrix A that each codeword must satisfy). The messages can take three different values: $*$ (erasure) or 0 or 1. Initially, each bit sends its value 0 or 1, or $*$ if erased, to each of the parity checks it is involved in. Check-to-bit and bit-to-check messages are then sent alternatively. If a check a receives non-erasure messages from all its bits but i , it sends to i the sum (modulo 2) of these messages; otherwise, the check a sends $*$ to i . If an erased bit i receives at least one non-erasure message from any of its checks but a , it sends it to a (if more than one, they are necessarily identical); otherwise, the bit i sends its value, 0 or 1, or $*$ if erased, to a . The algorithm stops after convergence of the iterations.

The (typical) cavity method, or BP density evolution, analyzes the outcome of this procedure in the limit where the codeword length N is infinite. It introduces η , the probability that a bit sends an erasure message to a check, and ζ the probability that a check sends an erasure message to a bit, both taken after BP has reached convergence. The *cavity equations* satisfied by these two probabilities,

$$\zeta = 1 - (1 - \eta)^{k-1}, \quad \eta = p\zeta^{\ell-1}, \quad (11)$$

characterize the fixed point of the BP density evolution (see Fig. 1).

Once BP has converged, bits receiving at least one non-erasure message are fixed to their correct value, as are the non-erased bits. When eliminated, along with the checks receiving no more than one erasure message, they leave the so-called *core*. The dimensions $M_c \times N_c$ of the associated residual matrix are, with high probability:

$$\begin{aligned} N_c &= p\zeta^\ell N + o(N), \\ M_c &= \frac{\ell}{k} [1 - (1 - \eta)^k - k\eta(1 - \eta)^{k-1}] N + o(N). \end{aligned} \quad (12)$$

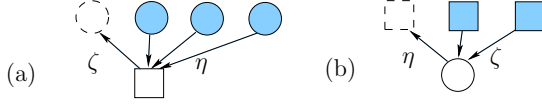


Figure 1: Illustration of the cavity equations (11), with $k = 4$ and $\ell = 3$. (a): a check node (square) sends an erasure message to a bit node (dashed circle) if at least one of its other variables sends an erasure message. (b): a bit node (circle) sends an erasure message to a check node (dashed square) if it has been erased and if all its other checks send an erasure message.

For $p < p_d(\ell, k)$, the only solution to (11) is $\zeta = 0, \eta = 0$, meaning that BP is able to decode the whole word with high probability. For $p > p_d$ however, BP gets stuck at some $\zeta > 0, \eta > 0$. In this case, it can be proved that the residual matrix has full-rank with high probability [9]. Therefore, the problem has exactly $2^{N_c - M_c}$ solutions if $N_c > M_c$, and one solution (the original codeword) otherwise. In this approach, the critical noise $p_c(\ell, k)$ is obtained from the condition $N_c = M_c$, and $\bar{s}$ is given by $\max(0, \bar{s}_{\text{cav}})$, with $\bar{s}_{\text{cav}} = \lim_{N \rightarrow \infty} (N_c - M_c)/N$.

The large deviation cavity method is built on the same ideas but incorporates a biased measure over the noise and code ensemble, as prescribed by Eq. (9). When we consider the value of a bit-to-check message as a function of its $(\ell - 1)(k - 1)$ “grandparents”, we also evaluate the “entropy shift” ΔS associated with the addition of the bit and its $\ell - 1$ checks, i.e. the difference between the numbers of columns and lines contributed by the bit and its checks to the residual matrix. Then the message is sent with a probability proportional to

$$(\mathbb{E} \xi^{2^{\Delta S}})^y. \quad (13)$$

For regular LDPC codes, we thus obtain for the potential

$$\begin{aligned} \psi(x, y) = \\ \log Z_\ell - \frac{\ell(k-1)}{k} \log [(1-\eta)^k + (1-(1-\eta)^k)2^{-xy}] \end{aligned} \quad (14)$$

with

$$Z_\ell = (\zeta 2^{-xy} + 1 - \zeta)^\ell - (\zeta 2^{-xy})^\ell + \zeta^\ell (p 2^x + 1 - p)^y 2^{-\ell xy} \quad (15)$$

and

$$\begin{aligned} \eta &= \zeta^{\ell-1} (p 2^x)^y 2^{-(\ell-1)xy} Z_{\ell-1}^{-1}, \\ \zeta &= 1 - (1-\eta)^{k-1}. \end{aligned} \quad (16)$$

Note that the entropy conjugated with x is not the “real” entropy s , but $s_{\text{cav}} = (N_c - M_c)/N$. When $x = 0$, the fixed point of the usual density evolution equations (11) is recovered, with $(1/y)\partial_x \psi(x=0, y)$ giving back $\bar{s}_{\text{cav}}$, the typical value.

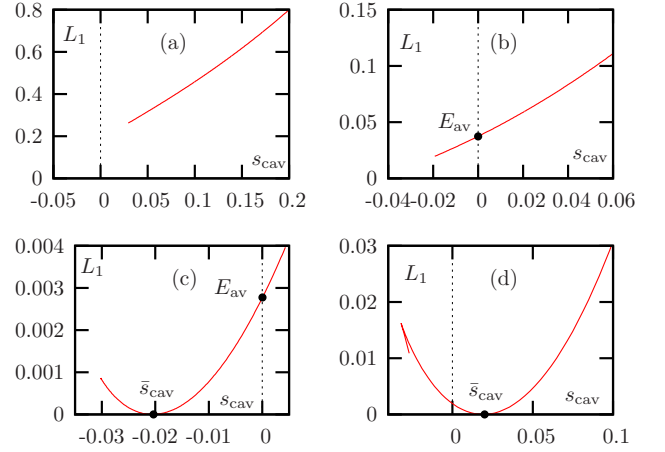


Figure 2: Average entropic rate function $L_1(s)$ as a function of the entropy density s_{cav} , for the regular LDPC code $\ell = 3, k = 6$ on the BEC with increasing values of p . The real entropy is actually $s = \max(0, s_{\text{cav}})$. (a): $p < p_{\text{1rsb}}$, no solution with $s = 0$; (b): $p_{\text{1rsb}} < p < p_d$, a solution with $s = 0$, but $\bar{s}$ is not defined; (c): $p_d < p < p_c$, $\bar{s} = 0$; (d): $p > p_c$, $\bar{s} > 0$ indicates that decoding typically fails.

IV. LDPC CODES

We first discuss average error exponents. The calculation of the average rate function $L_1(s)$ reveals four distinct regimes when the noise level p is varied, as illustrated and explained in Fig. 2. In particular, we find that the rate function $L_1(s)$ is no longer defined for $s = 0$ when p is too small ($p < p_{\text{1rsb}}$), which points to the inadequacy of our method in this low-noise regime.

Indeed, by retaining $s = 0$ as criterion for correct decoding, we assumed that an error implicates an exponential number of codewords. An error may however also be caused by the presence of one (or a few) isolated codeword(s). Estimating this probability requires an alternative, “energetic”, scheme, as opposed to the “entropic” scheme discussed so far³. Equations for the energetic average and typical error exponents can also be obtained from the large deviation cavity method [11], but their solutions are confined to a restricted interval $p > p_{\text{rs}}$, indicating again that the lowest noise levels are not appropriately described. The entropic and energetic exponents are found to cross at p_e , which corresponds to the so-called *critical rate* [1, 6]. We conjecture that the entropic exponent, as given by the above equations, is exact in the range $[p_e, p_c]$, while the energetic exponent (not presented here), which applies for $[p_{\text{rs}}, p_e]$, is only approximate.

³The energetic version of the cavity method is also referred to as “replica symmetric” in the physics literature, while the entropic version is known as “one-step replica symmetry breaking”.

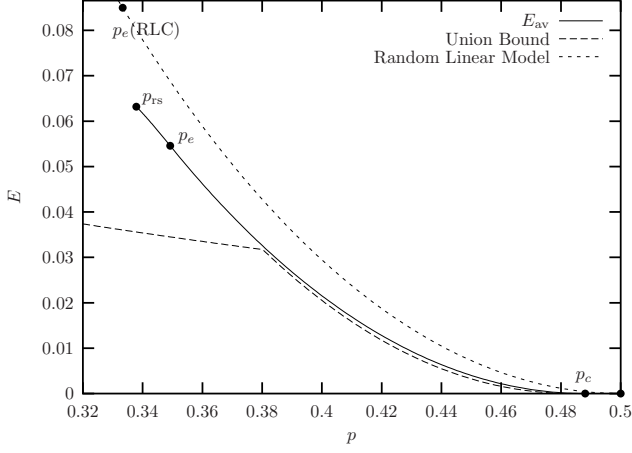


Figure 3: Average error exponent as a function of the noise level p of the BEC for the regular LDPC code ensemble with $k = 6$ and $\ell = 3$. Gallager's union bound and the random linear code limit (19) are also plotted for comparison.

(k, ℓ)	(4, 3)	(6, 3)
p_{rsb}	0.3252629709	0.2668568754
p_{rs}	0.5465748811	0.3378374641
p_e	0.6068720166	0.3491884902
p_d	0.6474256494	0.4294398144
p_c	0.7460097025	0.4881508842

Table 1: Thresholds p_{rsb} , p_{rs} , p_e , p_d and p_c (see text and Fig. 2) for two regular ensembles of LDPC codes.

Fig. 3 shows our predictions for the average exponent of the $\ell = 3$, $k = 6$ regular LDPC codes, with the two regimes represented; the same general picture holds for other regular or irregular ensembles (see also Table 1).

V. THE RANDOM LINEAR CODE LIMIT

This limit is obtained from regular codes with $k, \ell \rightarrow \infty$ and $R = 1 - \ell/k$ fixed, where the potential simplifies to:

$$\psi(x, y) = y \log(p2^x + 1 - p) + (R - 1)xy. \quad (17)$$

The trivial dependence of $\psi(x, y)$ with y implies that the two error exponents E_{av} and E_{typ} , as obtained from the entropic scheme, are identical. They are equal to the *volume bound* [2] $D(1 - R||p)$, where $D(x||y) = x \log(x/y) + (1 - x) \log((1 - x)/(1 - y))$ denotes the *Kullback-Leibler divergence*.

The intersection of the entropic and energetic average error exponents yields the threshold

$$p_e = \frac{1 - R}{1 + R}, \quad (18)$$

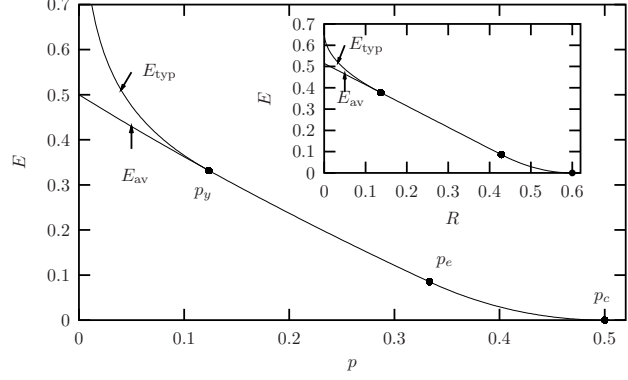


Figure 4: Average and typical error exponents of random linear codes on the BEC as a function of p , with $R = 1/2$ fixed. Inset: the same exponents as a function of R , with $p = 0.4$ fixed.

and we obtain for the average error exponent in the infinite connectivity limit:

$$E_{\text{av}}(\text{RLC}) = \begin{cases} 1 - R - \log(1 + p) & \text{if } p < p_e, \\ D(1 - R||p) & \text{if } p_e < p < p_c. \end{cases} \quad (19)$$

It coincides with the average error exponent of the *random linear code* (RLC) ensemble, where the $M \times N$ parity-check matrix is chosen at random with uniform probability among all possible parity-check matrices. Assuming that the inversion of the limits $N \rightarrow \infty$ and $k, \ell \rightarrow \infty$ is justified, we interpret this result as a validation of our approach (note that here, $p_{\text{rs}} = 0$).

The analysis of the typical error exponent in the energetic regime leads us to introduce an additional threshold,

$$p_y = \frac{\delta_{GV}(R)}{1 - \delta_{GV}(R)}, \quad (20)$$

where $\delta_{GV}(R)$, the minimal reduced distance of a typical linear code [1], is given by the smallest solution of $-\delta \log \delta - (1 - \delta) \log(1 - \delta) = 1 - R$. Below p_y , physical arguments [11] indicates that the typical error exponent must differ from the average one, with:

$$E_{\text{typ}}(\text{RLC}) = \begin{cases} -\delta_{GV}(R) \log p & \text{if } p < p_y, \\ E_{\text{av}}(\text{RLC}) & \text{if } p > p_y. \end{cases} \quad (21)$$

We are not aware of any previous report of this expression in the literature, but the fact that it matches the *union bound* suggests that it is exact. Fig. 4 presents the error exponents as a function of p for a fixed value of the rate $R = 1/2$.

The two thresholds p_e and p_y are presumably generic features of block codes, and are also found with random codes on the binary symmetric channel [1].

VI. DISCUSSION

Despite being one of the earliest and most basic topics in information theory, error exponents still retain today a number of unsolved issues. We advocated here a novel, thermodynamical, formulation of this problem. Using the cavity method from statistical mechanics, we worked out in this framework expressions for the average and typical error exponents of LDPC codes on the BEC. Our method provides an alternative to the replica method, applied to the BSC in [14], with the advantage of being based on explicit probabilistic assumptions. Our approach helps clarify the nature of the phase diagram, while the extension to the BEC allows for an analytical treatment.

While non rigorous, the cavity method aims at providing exact formulæ. Accordingly, our expressions are consistent with the various rigorous studies reported in the literature. The quest for rigorous proofs of formulæ obtained from the cavity method is currently an active field of mathematics [15]. Remarkably, predictions from the cavity method on the maximum-likelihood threshold p_c [4] could be turned into rigorous theorems [9]. This may inspire alternative derivations of our results.

Perhaps not too surprisingly, the entropic range $p_e < p < p_c$ where we conjecture our results to be exact also coincides with the limited interval for which the related problem of determining the reliability function of block codes has been solved so far. Extending our method to $p < p_e$, where we could obtain only approximate results (except in the infinite connectivity limit), remains a challenging open problem.

Using the same approach, we also analyzed the case of the binary symmetric channel, obtaining comparable results [11]. A more interesting extension would be to iterative decoding, such as BP. Although arguably quite academic, studying maximum-likelihood decoding, as we did, is nevertheless certainly an essential preliminary step.

ACKNOWLEDGMENTS

It is a pleasure to thank Stefano Ciliberti, Marc Mézard and Lenka Zdeborová for their critical reading. The work of T.M. was supported in part by the EC through the network MTR 2002-00319 ‘STIPCO’ and the FP6 IST consortium ‘EVERGROW’. O.R. is a fellow of the Human Frontier Science Program.

REFERENCES

- [1] A. Barg and G. D. Forney Jr., “Random codes : minimum distances and error exponents,” *IEEE Trans. Inform. Theory*, 48:2568–2573, 2002.
- [2] E. R. Berlekamp, “The performance of block codes,” *Notices of the AMS*, pages 17–22, January 2002.
- [3] F. den Hollander, *Large deviations*, Fields Institute Monographs 14. American Mathematical Society, Providence RI, 2000.
- [4] S. Franz, M. Leone, A. Montanari, and F. Ricci-Tersenghi, “The dynamic phase transition for decoding algorithms,” *Phys. Rev. E*, 66:046120, 2002.
- [5] R. G. Gallager, “Low-density parity check codes,” *IRE Trans. Inf. Theory*, IT-8:21, 1962.
- [6] R. G. Gallager, *Information theory and reliable communication*, John Wiley and Sons, New York, 1968.
- [7] Y. Kabashima and D. Saad, “Statistical mechanics of low-density parity-check codes,” *J. Phys. A: Math. Gen.*, 37:R1–R43, 2004.
- [8] M. Luby, M. Mitzenmacher, A. Shokrollahi, and D. Spielman, “Efficient erasure correcting codes,” *IEEE Trans. Inform. Theory*, vol. 47, 569–584, Feb. 2001.
- [9] C. Measson, A. Montanari, T. Richardson, and R. Urbanke, “Life above threshold: from list decoding to area theorem and MSE,” In *Proc. ITW*, San Antonio, USA, October 2004.
- [10] M. Mézard and G. Parisi. “The Bethe lattice spin glass revisited,” *Eur. Phys. J. B*, 20:217, 2001.
- [11] T. Mora and O. Rivoire, 2006. In preparation.
- [12] O. Rivoire. “The cavity method for large deviations,” *J. Stat. Mech.*, P07004, 2005.
- [13] D. Ruelle. *Thermodynamic formalism*, Cambridge Math. Library, 2nd Ed, 2004.
- [14] N. S. Skantzos, J. van Mourik, D. Saad, and Y. Kabashima, “Average and reliability error exponents in low-density parity-check codes,” *J. Phys. A*, 36:11131–11141, 2003.
- [15] M. Talagrand, *Spin glasses : a challenge for mathematicians. Cavity and mean field models*, Springer-Verlag, New-York, 2003.

“Statistical mechanics of error exponents for
error-correcting codes”

Phys. Rev. E **74**, 056110 (2006)

Statistical mechanics of error exponents for error-correcting codes

Thierry Mora

Laboratoire de Physique Théorique et Modèles Statistiques, Bât. 100, Université Paris-Sud, F-91405 Orsay, France

Olivier Rivoire

Laboratory of Living Matter, The Rockefeller University, 1230 York Avenue, Box 34, New York, New York 10021, USA

(Received 27 June 2006; published 15 November 2006)

Error exponents characterize the exponential decay, when increasing message length, of the probability of error of many error-correcting codes. To tackle the long-standing problem of computing them exactly, we introduce a general, thermodynamic, formalism that we illustrate with maximum-likelihood decoding of low-density parity-check codes on the binary erasure channel and the binary symmetric channel. In this formalism, we apply the cavity method for large deviations to derive expressions for both the average and typical error exponents, which differ by the procedure used to select the codes from specified ensembles. When decreasing the noise intensity, we find that two phase transitions take place, at two different levels: a glass to ferromagnetic transition in the space of codewords and a paramagnetic to glass transition in the space of codes.

DOI: [10.1103/PhysRevE.74.056110](https://doi.org/10.1103/PhysRevE.74.056110)

PACS number(s): 89.90+n, 89.70+c, 05.50+q

I. INTRODUCTION

Communicating information requires a physical channel whose inherent noise impairs the transmitted signals. Reliability can be improved by adding redundancy to the messages, thus allowing the receiver to correct the effects of the noise. This procedure has the drawbacks of increasing the cost of generating and sending the messages and of decreasing the speed of transmission. At first sight, better accuracy seems achievable only at the expense of lesser efficiency. Remarkably, Shannon showed that, in the limit of infinite-length messages, error-free communication is possible using only limited redundancy [1]. His proof of principle has triggered many efforts to construct actual error-correcting schemes that would approach the theoretical bounds. A renewal of interest in the subject has taken place during the last ten years, as new error-correcting codes were finally discovered [2], or rediscovered [3], which showed practical performances close to Shannon's bounds.

In this paper, we analyze a major family of such codes, the low-density parity-check (LDPC) codes, also known as Gallager codes, from the name of their inventor [4]. Our focus is on the characterization of rare decoding errors, in situations where most realizations of the noise are accurately corrected. Error-free communication, as guaranteed by Shannon's theorem, indeed results from a law of large number and is achieved only with infinite-length messages. Accordingly, any error-correcting scheme acting on finite-length messages has a nonzero error probability, which generically vanishes exponentially with the message length. Such error probabilities are described by *error exponents*, giving their rate of exponential decay. Two kinds of error exponents are usually distinguished: *average* error exponents, where the average is taken over an ensemble of codes, and *typical* error exponents, where the codes are typical elements of their ensemble.

The study of error exponents attracted early on considerable attention in the information theory community, but exact expressions have turned out to be particularly difficult to derive (see, e.g., [5] and [6] for concise and nontechnical

reviews with entries in the literature). Exact asymptotic results are known in the limit of the so-called *random linear model* [7] (presented in Appendix B), but only loose bounds (presented in Appendix C) have been established for more general codes. Recently, a systematic finite-length analysis of LDPC codes under iterative decoding was carried out for the binary erasure channel (BEC) [8,9], yielding exact, yet non-explicit, formulas for the average error probability. Up to now, little has, however, been known of the error probability under maximum-likelihood decoding, except for the work of [10] dealing with the binary symmetric channel (BSC).

We address here the problem of computing error exponents of LDPC codes under maximum-likelihood decoding, over both the BEC and BSC (all the necessary definitions are recalled below). We adopt a statistical physics point of view, which exploits the well-established [11] mapping between error-correcting codes and spin glasses [12]. A thermodynamic formalism is introduced where error exponents are expressed as large deviation functions [13], which we compute by means of the extension of the cavity method [14] proposed in [15]. This approach offers an alternative to the related replica method employed in [10] and allows us to address both average and typical error exponents. We thus obtain an interesting phase diagram, with two very distinct phase transitions occurring when the intensity of the noise in the channels is varied.

A brief summary of our results can be found in [16]. We present in what follows a much more detailed account of our approach. In a first part, we define LDPC codes, recall their mapping to some models of spin glasses and optimization problems, and give a general overview of our thermodynamic (large deviation) formalism. The two subsequent parts apply this framework to the analysis of LDPC codes over the BEC and BSC, respectively. We sum up our results in a conclusion where we also point out some open questions. Most of the technical calculations are relegated to the Appendixes, which also contain a detailed discussion of the limiting case of random linear codes.

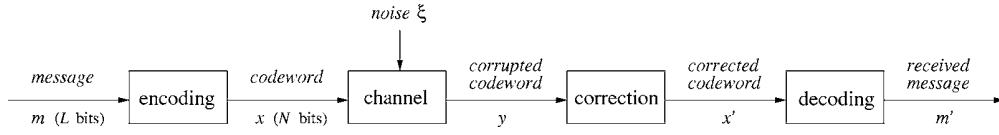


FIG. 1. Error correction scheme. A message m composed of L bits, $m \in \{0,1\}^L$, is first encoded in a codeword of longer size N with $R=L/N < 1$, defining the rate of the code. The noise ξ of the channel corrupts the transmitted codeword which becomes y (see Fig. 2 for examples of channels). This output is generically not a codeword, and the correction consists in inferring the most probable codeword to which it comes from. Finally, the inferred codeword x' is converted back into its corresponding message m' . The communication is successful if $m'=m$.

II. ERROR-CORRECTING CODES AND THE LARGE DEVIATION FORMALISM

A. Error-correcting codes

Error-correcting codes are based on the idea that adding sufficient redundancy to the messages can allow the receiver to reconstruct them, even if they have been partially corrupted by the noisy channel [17]. A schematic view of how these codes operate is presented in Fig. 1. Given a message composed of L bits, an encoding map $\{0,1\}^L \rightarrow \{0,1\}^N$ first introduces redundancy by converting the L bits of the message into a longer sequence of N bits, called a *codeword*. The ratio $R \equiv L/N$ defines the *rate* of the code and should ideally be as large as possible to reduce communication costs, yet small enough to allow for corrections. Corrections are implemented downstream the noisy channel and specified by a decoding map $\{0,1\}^N \rightarrow \{0,1\}^L$ whose purpose is to reconstruct the original message from the received corrupted codeword. Decoding is composed of two steps: first, the most probable codeword is inferred, and second, it is converted into its corresponding message.

In this scheme, messages and codewords are related by the one-to-one encoding map, and translating messages into codewords or conversely is relatively straightforward. The computationally most demanding part is concentrated on inferring the most probable codeword sent, given the corrupted codeword received. In what follows, we shall focus exclusively on this problem, which requires manipulating only codewords.

B. Communication channels

Formally, a noisy channel is characterized by a transition probability $Q(\mathbf{y}|\mathbf{x})$ giving the probability for its output to be $\mathbf{y}$ given that its input was $\mathbf{x}$. For the sake of simplicity, we confine ourselves to *memoryless* channels where the noise affects each bit independently of the others—i.e., $Q(\mathbf{y}|\mathbf{x}) = \prod_{i=1}^N Q(y_i|x_i)$ with $Q(y_i|x_i)$ independent of i .

We shall consider more specifically two examples of memoryless channels. The first one is the *binary erasure channel* where a bit is erased with probability p —that is, $Q(*|x)=p$ and $Q(x|x)=1-p$ where $*$ represents an erased bit (see Fig. 2). The second is the *binary symmetric channel* where a bit is flipped with probability p —that is, $Q(0|1)=Q(1|0)=p$ and $Q(0|0)=Q(1|1)=1-p$ (see Fig. 2).

C. LDPC codes and code ensembles

Shannon first formalized the problem of error correction and determined the lowest achievable rate R allowing error-

free correction [1]. He found a general expression for this limit, called the *channel capacity*, which depends only on the nature of the channel and takes the form $C_{\text{BEC}}(p)=1-p$ and $C_{\text{BSC}}(p)=1-p \ln p - (1-p) \ln(1-p)$ for the BEC and BSC, respectively. Shannon's proof for the existence of codes achieving the channel capacity was nonconstructive and his analysis restricted to the limit of infinitely long messages, $L \rightarrow \infty$. Among the various families of codes proposed to practically perform error correction, one of the most promising is the family of *low-density parity-check* codes [4].

A LDPC code is defined by a sparse matrix A where “sparse” means that A is mostly composed of 0's, with otherwise a few 1's. The *parity-check matrix* A has size $M \times N$ with $M=N-L$ and is associated with a *generator matrix* G of size $L \times N$ such that $GA=0$ (see, e.g., [3] for explicit constructions); the encoding map is taken to be the linear map $x=Gm$ and the rate of the code is $R=L/N=1-M/N$. By construction, an N -bit codeword x satisfies the M parity-check equations $Ax=0$, or, in other words, the set of codewords is the kernel of A . The parity-check matrix A is usually represented graphically by a *factor graph*, as in Fig. 3: the columns of A are associated with *check nodes* labeled with $a \in \{1, \dots, M\}$ and represented by squares, and the lines of A are associated with *variable nodes* labeled with $i \in \{1, \dots, N\}$ and represented by circles. A nonzero element of the matrix A such as $A_{ia}=1$ appears as a link between the variable node i and the check node a .

A particularly powerful approach for analyzing error-correcting codes is the probabilistic method where, instead of considering a single code, one studies an *ensemble* of codes. With LDPC codes, code ensembles correspond to sets of matrices or, equivalently, sets of factor graphs. A popular choice is to consider the ensemble of factor graphs with given connectivities c_k and v_ℓ , which is the set of factor graphs having $c_k M$ check nodes with connectivity k and $v_\ell N$ variable nodes with connectivity ℓ , where $\sum_k c_k = \sum_\ell v_\ell = 1$. A convenient representation is by means of the generating

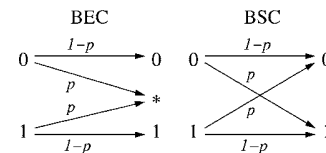


FIG. 2. Communication channels. On the left the BEC (binary erasure channel) erases a bit with probability p and leaves it unchanged with probability $1-p$. On the right the BSC (binary symmetric channel) flips a bit with probability p and leaves it unchanged with probability $1-p$.

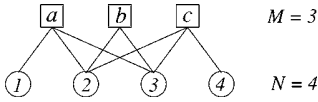


FIG. 3. Factor graph (Tanner graph [18]). The circles represent the variable nodes, associated with the N bits $\{x_i\}$, and the squares represent the M parity check. In the example given, the constraints read: (a) $x_1 + x_2 + x_3 = 0$, (b) $x_2 + x_3 = 0$, and (c) $x_2 + x_3 + x_4 = 0$ (modulo 2).

functions $c(x) = \sum_k c_k x^k$ and $v_\ell = \sum_\ell v_\ell x^\ell$; these notations allow one, for instance, to write the mean connectivities as $\langle k \rangle = c'(1)$ and $\langle \ell \rangle = v'(1)$. Due to their simplicity, particular attention will be devoted to *regular* codes, whose check nodes have all same degree k and variable nodes same degree ℓ , corresponding to $c_{k'} = \delta_{k,k'}$ and $v_{\ell'} = \delta_{\ell,\ell'}$ or, equivalently, $c(x) = x^k$ and $v(x) = x^\ell$.

The mathematical fact underlying the probabilistic method is the phenomenon of *measure concentration* which occurs in the limit where $N \rightarrow \infty$ and $M \rightarrow \infty$ with fixed ratio $\alpha = M/N$: in this limit, many properties are shared by *almost all* elements of the ensemble (i.e., all but a subset of measure zero). As a consequence, by studying average properties over an ensemble, one actually has access to properties of typical elements of this ensemble. This fact is one of the building blocks of random graph theory [19] and is also central to the physics of disordered systems where it is known as the *self-averaging property* [20].

While the factor graph representation makes obvious the connection between LDPC codes and random graph theory, it will also turn particularly fruitful to exploit the close ties of LDPC codes with both optimization problems [21] and spin-glass systems [20]. LDPC codes are indeed intimately related to a class of combinatorial optimization problems known as XORSAT problems where, given a sparse matrix A and a vector τ , one is to find solutions σ to the equation $A\sigma = \tau$. Although algorithmically relatively simple (Gauss method provides an answer in a time polynomial in the size of the matrix), XORSAT problems share many common features with notably more difficult, NP-complete [21], problems such as K -SAT. A recent physical approach to XORSAT problems makes use of their formal equivalence with a class of spin-glass systems known as p -spin models [22–24]. We shall follow this line of investigation and apply the *cavity method* [14,25] from spin-glass theory to analyze LDPC codes. We note that alternative, sometimes equivalent, physical approaches have previously been applied to LDPC codes; we refer the reader to [26] for a review of the subject.

The distinctive feature of XORSAT at the root of its computational simplicity is the presence of an underlying group symmetry that relates all solutions. In the context of LDPC codes, it corresponds to the fact that the set of codewords is the kernel of the parity-check matrix A ; we shall refer to the XORSAT problem $A\sigma = 0$ whose solutions define the set of codewords as the *encoding constraint satisfaction problem* (CSP) of the LDPC code with check matrix A . The group symmetry has a number of interesting consequences which will crucially simplify the analysis.

Most of the interest in LDPC codes stems from the possibility to decode them using efficient, iterative algorithms

(described in Sec. III A 3). Unless otherwise stated, we shall, however, be here concerned with the theoretically simpler, yet computationally much more demanding, *maximum-likelihood decoding* procedure. It consists in systematically decoding a received message to the most probable codeword (a task that iterative algorithms are in some cases unable to perform, as recalled in Sec. III A 3).

Finally, it is interesting to note that in the limit where $\langle k \rangle, \langle \ell \rangle \rightarrow \infty$ with fixed ratio, LDPC codes define the *random linear model* (RLM) whose typical elements have been shown by Shannon to achieve the channel capacity. This particular limit, where many quantities can be computed by invoking only elementary combinatorial arguments, is discussed in detail in Appendix B.

D. Typical properties and phase transitions

The performance of a particular code over a given channel is measured by its error probability—i.e., the probability that it fails to correctly decode a corrupted codeword. More precisely, if $d(\mathbf{y})$ denotes the inferred codeword when $\mathbf{x}$ is sent and $\mathbf{y}$ received, one defines the *block error probability* for $\mathbf{x}$ as

$$P_N^{(B)}(\mathbf{x}) = \sum_{\mathbf{y}} Q(\mathbf{y}|\mathbf{x}) 1_{d(\mathbf{y}) \neq \mathbf{x}} \quad (1)$$

and the average block error probability as

$$P_N^{(B)} = \mathbb{E}_{\mathbf{x}}[P_N^{(B)}(\mathbf{x})], \quad (2)$$

where $\mathbb{E}_{\mathbf{x}}$ denotes the expectation (average) over the set of codewords. With LDPC codes, this average is trivial since, due to the group symmetry, all codewords are equivalent, and $P_N^{(B)}(\mathbf{x})$ is independent of $\mathbf{x}$.

The concentration phenomenon alluded to above means here that $P_N^{(B)} \rightarrow p_B$ with $N \rightarrow \infty$ within a given code ensemble defined by generating functions $c(x)$ and $v(x)$. As the level of the noise p is increased, a phase transition is generically observed: a critical value p_c exists above which error-free correction is no longer possible ($p_B = 0$ for $p < p_c$ and $p_B = 1$ for $p > p_c$). The formalism to be presented in the next sections will yield in particular the value of p_c for given code ensembles and channels. Obviously, the presence of this phase transition indicates that, when using a channel with noise level p , one should choose a code from an ensemble for which $p < p_c$. The phase transition is, however, occurring only in the limit of infinite codewords (thermodynamic limit) whereas practical coding inevitably deals with finite N . This leads to the fact that the block error probability is not exactly zero, even in the regime $p < p_c$.

For a given code of finite but large block-length N , error can thus be caused by rare, atypical, realizations of the noise. Similarly, when picking a code at random from a code ensemble of finite size, one can observe properties differing from the typical properties predicted by the law of large numbers. We show in what follows how these two atypical features induced by finite-size effects can be analyzed in a common framework.

TABLE I. The analogy with spin glasses or, more generally, the statistical physics of disordered system with quenched disorder.

	Spin glass	Average	Typical	Multistep, step 1	Multistep, step 2
Disorder	Couplings J_{ij}		Typical codes $\mathcal{C}^0$	Codes $\mathcal{C}$ at y	
Configurations	Spins $\{\sigma_i\}_i$	Noise + codes $(\xi, \mathcal{C})$	Noise ξ	Noise ξ	Codes $\mathcal{C}$
Observable	$E = \sum_{ij} J_{ij} \sigma_i \sigma_j$	$S_N(\xi, \mathcal{C})$	$S_N(\xi, \mathcal{C}^0)$	$S_N(\xi, \mathcal{C})$	$L_{\mathcal{C}}(s)$
Entropy	$s(e = E/N)$	$L_0(s = S_N/N)$	$L(s = S_N/N)$	$\mathcal{L}(\phi, x)$	
Temperature ⁻¹	$\beta = \partial_e s$	$x = \partial_s L_1$	$x = \partial_s L_0$	$x = \partial_s L$	$y = \partial_\phi \mathcal{L}$
Potential	$\beta f = \beta e - s$	$\phi_1 = xs - L_1$	$\phi_0 = xs - L_0$	$\phi = xs - L$	$\psi = y\phi - \mathcal{L}$

E. Large deviations

At this stage, it is useful to make explicit the three different levels of statistics involved in the analysis of error-correcting codes: (i) statistics over the codes $\mathcal{C}$ in a defined code ensemble $\mathcal{C}$, (ii) statistics over the set of transmitted codewords $\mathbf{x}$ of a particular code, and (iii) statistics over the noise ξ of the channel, with a specified p . For given $\mathcal{C}$, $\mathbf{x}$, and ξ , a fourth level of statistics is involved in the decoding process, over the possible codewords $y \in \{0, 1\}^N$ from which the received corrupted codeword originates. The group structure of the set of codewords of LDPC codes makes level (ii) trivial since all codewords are in fact equivalent (isomorphic). We will consequently ignore it and address only levels (i) and (iii).

The problem of evaluating the probability that, due to finite-size effects, a property differs from the typical case belongs to *large deviation* theory [13]. To give here a general presentation of the concepts and methods to be used, we assume that the success of the decoding is measured by a function $S_N(\xi, \mathcal{C})$ extensive in N and such that $S_N(\xi, \mathcal{C}) \leq 0$ if the code $\mathcal{C}$ correctly decodes a message subject to noise ξ and $S_N(\xi, \mathcal{C}) > 0$ otherwise; in the next sections, we will show explicitly how such an observable can be defined with LDPC codes, for both the BEC and BSC channels. In terms of S_N , the decoding phase transition takes the following form: in the limit $N \rightarrow \infty$, the distribution of the density S_N/N concentrates around a typical value $s_{\text{typ}}(p)$ which verifies $s_{\text{typ}}(p) \leq 0$ if $p < p_c$, and $s_{\text{typ}}(p) > 0$ if $p > p_c$, where p denotes as before the level of noise of the channel (see Fig. 2 for examples).

For typical codes in their ensemble, denoted $\mathcal{C}^0$, we describe large deviations of S_N with respect to the noise ξ by a *rate function* $L_0(s)$ such that the probability to observe $S_N(\xi, \mathcal{C}^0)/N = s$ satisfies

$$P_N[\xi; S_N(\xi, \mathcal{C}^0)/N = s] \asymp e^{-NL_0(s)}. \quad (3)$$

Here the symbol $a_N \asymp b_N$ refers to an exponential equivalence, $\ln a_N / \ln b_N \rightarrow 1$ as $N \rightarrow \infty$. Viewed as a function of the noise level p , the rate function $E_{\text{typ}}(p) = L_0(s=0)$ is known in the coding literature as the *typical error exponent* [5]. The exponential decay with N of atypical properties is quite generic when dealing with large deviations, but this scaling is not necessarily ensured, as discussed in more detail in Appendix A. In the thermodynamic formalism that we shall

adopt, rate functions are computed by introducing a potential $\Phi_{\mathcal{C}}(x)$ defined by

$$\Phi_{\mathcal{C}}(x) = \ln(\mathbb{E}_{\xi}[e^{xS_N(\xi, \mathcal{C})}]). \quad (4)$$

In the limit $N \rightarrow \infty$ limit, the density $\Phi_{\mathcal{C}}(x)/N$ tends to a typical value $\phi_0(x)$, which is related to the rate function $L_0(s)$ by

$$e^{N\phi_0(x)} \asymp \int ds e^{N[xs - L_0(s)]}. \quad (5)$$

Equivalently, by taking the saddle point,

$$\phi_0(x) = xs - L_0(s), \quad x = \partial_s L_0(s). \quad (6)$$

The rate function $L_0(s)$ can thus be reconstructed from $\phi_0(x)$ by inverting the Legendre transformation,

$$L_0(s) = sx - \phi_0(x), \quad s = \partial_x \phi_0(x). \quad (7)$$

The analogy with the usual thermodynamics is summarized in Table I.

From a theoretical perspective, it is simpler to make an average over the codes and compute the rate function $L_1(s)$ defined as

$$P_N[\xi, \mathcal{C}; S_N(\xi, \mathcal{C})/N = s] \asymp e^{-NL_1(s)}. \quad (8)$$

This procedure yields the so-called *average error exponent* $E_{\text{av}} = L_1(s=0)$. In the thermodynamical formalism, $L_1(s)$ is conjugated to the potential $\phi_1(x)$ satisfying

$$e^{N\phi_1(x)} = \mathbb{E}_{(\xi, \mathcal{C})}[e^{xS_N(\xi, \mathcal{C})}] = \int ds e^{N[xs - L_1(s)]}. \quad (9)$$

The two rate functions $L_0(s)$ and $L_1(s)$ may differ, meaning that the average exponent can be associated with atypical codes. Such atypical codes correspond themselves to large deviations of the potential $\Phi_{\mathcal{C}}(x)$. For fixed values of x , we define a rate function $\mathcal{L}(\phi, x)$ as

$$P_N[\mathcal{C}; \Phi_{\mathcal{C}}(x)/N = \phi] \asymp e^{-N\mathcal{L}(\phi, x)}. \quad (10)$$

In a thermodynamic formalism, $\mathcal{L}(\phi, x)$ is again associated with a potential $\psi(x, y)$ defined by

TABLE II. Analogy with the replica approach of spin glasses. The replica-symmetric method prescribes that the typical partition function Z_0 of a disordered system is given by $Z_0 \sim \mathbb{E}[Z_N^n]^{1/n}$ with $n \rightarrow 0$ or, more precisely, if $\Lambda_N = \ln Z_N$, the typical value of $\lambda = \Lambda_N/N$ is $\lambda_0 = \lim_{n \rightarrow 0} \lim_{N \rightarrow \infty} (1/Nn) \ln \mathbb{E}[e^{n\Lambda_N}]$. This is mathematically justified by the Gärtner-Ellis theorem which moreover provides a rigorous basis for the interpretation of nonzero values of n in terms of large deviations, as discussed in the text. According to this theorem, if the function $\phi(x) = \lim_{N \rightarrow \infty} (1/N) \ln \mathbb{E}[e^{x\Lambda_N}]$ exists and is regular enough (see, e.g., [13] for a rigorous presentation), then a large deviation principle holds for λ with a rate function being the Legendre transform of $\phi(x)$; if we assume the functions differentiable, $L(\lambda) = \lambda x - \phi(x)$ with $\lambda = \partial_x \phi(x)$. As a corollary of this theorem, the typical value λ_0 , which by definition satisfies $L(\lambda_0) = 0$ and $x = \partial_\lambda L(\lambda_0) = 0$, is given by $\lambda_0 = \partial_x \phi(x=0) = \lim_{x \rightarrow 0} [\phi(x)/x] (x=0)$, as predicted by the replica method. Note also that $n=1$, with $Z_1 = \mathbb{E}[Z_N]$, corresponds to the so-called annealed approximation.

Replica (symmetric) theory of spin glasses	Multistep large deviations for LDPC codes
Hamiltonian $H_L[\sigma] = \sum_{ij} J_{ij} \sigma_i \sigma_j$	$S_N(\xi, C)$
Disorder $\{J_{ij}\}_{ij}$	Codes C
Configurations $\{\sigma_i\}_i$	Noise ξ
Number of replicas n	Temperature $^{-1}$ y
Physical temperature $^{-1}$ β	Temperature $^{-1}$ x
Annealed approximation $n=1$	Average codes $y=1$
Quenched computation $n \rightarrow 0$	Typical codes $y \rightarrow 0$

$$e^{N\psi(x,y)} = \mathbb{E}_C[\mathbb{E}_\xi[e^{xS_N(\xi,C)}]]^y = \mathbb{E}_C[e^{y\Phi_C(x)}] = \int d\phi e^{N[y\phi - \mathcal{L}(\phi,x)]}. \quad (11)$$

We refer to this hierarchical embedding of large deviations as a *multistep large deviation* structure [15], a term meant to reflect the formal equivalence with the multistep replica symmetry breaking scenario developed for spin glasses [20] (see Table II). In the limit $N \rightarrow \infty$ where the integral is dominated by its saddle point we obtain the Legendre transformation

$$\psi(x,y) = y\phi - \mathcal{L}(\phi,x), \quad y = \partial_\phi \mathcal{L}(\phi,x). \quad (12)$$

Within this extended framework, we recover the average case by taking $y=1$. Indeed, from the definitions (9) of $\phi_1(x)$ and (11) of $\psi(x,y)$ it follows that

$$e^{N\psi(x,1)} = \mathbb{E}_C[\mathbb{E}_\xi[e^{xS_N(\xi,C)}]] = \mathbb{E}_{(\xi,C)}[e^{xS_N(\xi,C)}] = e^{N\phi_1(x)}, \quad (13)$$

that is,

$$\psi(x,y=1) = \phi_1(x). \quad (14)$$

This average case differs in general from the typical case which corresponds to $y=0$. Indeed, by definition [see Eq. (10)], typical codes are associated with the potential ϕ_0 minimizing $\mathcal{L}(\phi,x)$, with $\mathcal{L}(\phi_0,x)=0$, yielding $y = \partial_\phi \mathcal{L} = 0$. Note that the potential ϕ_0 is related to $\psi(x,y)$ by $\phi_0(x) = \lim_{y \rightarrow 0} (1/y) \psi(x,y)$, which can also be viewed as a corol-

lary of Gärtner-Ellis theorem [13], best known in statistical physics as the replica trick [20] (see Table II). In the language of the replica method, the average case ($y=1$) and the typical case ($y=0$) are, respectively, referred to as the annealed and quenched computations.

The previous discussion assumed that the potentials were analytical functions of their parameters x and y , but this may not be the case, and we will find that phase transitions can occur when these temperatures are varied. In such cases, taking naively the limit $y \rightarrow 0$ leads to erroneous results. We will discuss how to overcome such difficulties when encountering them.

III. LDPC CODES OVER THE BEC

We now proceed to illustrate our formalism with LDPC codes over the binary erasure channel. We start with rederiving the typical phase diagram by means of the cavity method, a slightly different approach than the replica method originally used in [27]. This sets the stage for the analysis of error exponents that follows.

A. Typical phase diagram

1. Formulation

Consider a LDPC code C with parity-check matrix A ; its *encoding CSP* (the constraint satisfaction problem whose SAT assignments define the codewords) has cost function

$$H_C[\sigma] = \sum_{a=1}^M E_a[\sigma], \quad \text{with } E_a[\sigma] = \sum_{i=1}^N A_{ai} \sigma_i \pmod{2}. \quad (15)$$

Since $E_a[\sigma] \in \{0,1\}$, the cost function $H_C[\sigma]$ counts the number of constraints violated by the assignment $\sigma = \{\sigma_i\}_{i=1,\dots,N}$ (where $\sigma_i \in \{0,1\}$). When a codeword σ^* , satisfying $H_C[\sigma^*] = 0$, goes through a BEC, each of its bits σ_i has probability p to be erased. A given realization of the noise can be characterized by a vector $\xi = (\xi_1, \dots, \xi_N)$ with $\xi_i = 1$ implying that the bit σ_i^* is lost and $\xi_i = 0$ that it is unaffected. If we denote by $\mathcal{E}$ the set of indices i for which $\xi_i = 1$ (erased bits), the decoding task consists in reconstructing $\{\sigma_i^*\}_{i \in \mathcal{E}}$ from the received bits $\{\sigma_i^*\}_{i \notin \mathcal{E}}$ and knowledge of the encoding CSP H_C . This decoding problem defines a new constraint satisfaction problem, the *decoding CSP*, obtained from the encoding CSP by fixing the values of the noncorrupted bits. More explicitly, the decoding CSP has cost function $H_C^{(\xi)}[\sigma^{(\xi)}] = \sum_a E_a^{(\xi)}[\sigma^{(\xi)}]$ where $\sigma^{(\xi)} = \{\sigma_i\}_{i \in \mathcal{E}}$ and

$$E_a^{(\xi)}[\sigma^{(\xi)}] = \sum_{i \in \mathcal{E}} A_{ai} \sigma_i + \sum_{i \notin \mathcal{E}} A_{ai} \sigma_i^* \pmod{2}. \quad (16)$$

Decoding is possible if and only if $\{\sigma_i^*\}_{i \in \mathcal{E}}$ is the only SAT assignment of the decoding CSP.

If $\mathcal{N}_N(\xi, C)$ denotes the number of solutions of the decoding CSP, $S_N(\xi, C)$ can be taken as $S_N(\xi, C) \equiv \ln \mathcal{N}_N(\xi, C)$. This entropy fulfills the desired properties: namely, $S_N(\xi, C) \leq 0$ if decoding is successful, and $S_N(\xi, C) > 0$ otherwise.

The particularity of LDPC codes compared to other error-correcting codes is that the decoding CSP has same form as the encoding CSP (both are XORSAT problems). As a consequence, the $\mathbb{Z}_2$ symmetry of the group of codewords is always preserved, at variance with what happens in other CSP's where fixing variables breaks a symmetry. The BEC is also particular compared with other channels, since the set $\mathcal{E}$ of corrupted bits is known to the receiver (this will not be the case with the BSC, where identifying the corrupted bits is part of the decoding problem). This entails that bits can only be fixed to their correct value.

2. Cavity approach

Before considering large deviations, it is instructive to recall the typical results—i.e., the values taken by $S_N(\xi, C^0)$ when C^0 is a typical code from a given ensemble specified by $c(x)$ and $v(x)$, and ξ a typical realization of the noise from the probability distribution specified by p . We resort here to the *cavity method at zero temperature* [14], whose validity is based on the treelike structure of the factor graphs associated with typical LDPC codes. The essentially equivalent replica method has been used in the past: in [28], $S_N(\xi, C)$ is thus obtained by first computing a free energy with the replica method and then taking the zero-temperature limit to obtain $S_N(\xi, C)$, viewed as the entropy of the zero-energy ground states.

The approach we follow here, which corresponds to a particular implementation of the entropic cavity method presented in [29], has several advantages over the replica approach: it involves neither a zero-replica limit nor a zero-temperature limit, it emphasizes the specificities of LDPC codes associated with the underlying $\mathbb{Z}_2$ symmetry, and it naturally connects to the algorithmic analysis of single codes. In the common language of the replica and cavity methods, the calculation to be done is coined *one-step replica symmetry breaking* (1RSB) and the entropy $s = S_N/N$ is referred to as a *complexity*. This is reflected in what follows by the fact that we strictly restrict to SAT assignments and assume that all constraints are satisfied (the reweighting parameter μ , as denoted in [25], is here infinite, $\mu = \infty$). This 1RSB approach is known to exactly describe XORSAT problems [23,24].

Let $P_i(\sigma_i)$ be the probability, taken over the set of solutions of the decoding CSP, that the bit i assumes the value $\sigma_i \in \{0, 1\}$. Due to the preservation of the $\mathbb{Z}_2$ symmetry, no bit can be nontrivially biased: either it is fixed to 0 or 1, corresponding to $P_i = \delta_0$ and $P_i = \delta_1$, respectively, or it is completely free, corresponding to $P_i = (\delta_0 + \delta_1)/2$, where we denote $\delta_\tau(\sigma) = \delta_{\tau, \sigma}$. In technical terms, the evanescent fields that are generically required to compute entropies in CSP [29] have here a trivial distribution, thus explaining that they can be safely ignored, as was done in [28].

Let ν be the probability, taken over the N nodes of a typical factor graph, that a bit i is free—i.e., that $P_i = (\delta_0 + \delta_1)/2$. Since a free node has equal probability to be 0 or 1, its contribution to the entropy is $\ln 2$ and the mean entropic contribution per node is $\nu \ln 2$. This value is, however, only an upper bound (known as the annealed, or first moment,

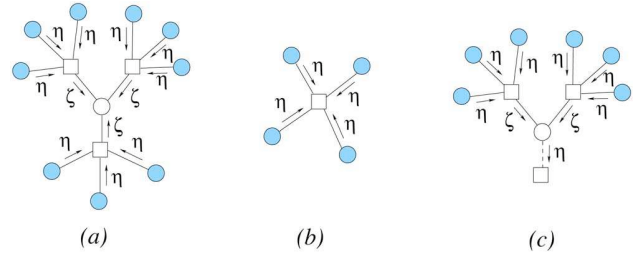


FIG. 4. (Color online) Illustration of cavity fields: (a) addition of a variable node, (b) addition of a parity check, and (c) cavity iteration.

bound) on the entropy density $s = S_N/N$ that we wish to calculate. In fact, it holds only if the bits are independent: indeed, two bits may both be free but, by fixing one, the second may be constrained to a unique value, in which case the joint entropic contribution of the two nodes is $\ln 2$ and not $2 \ln 2$. The correct expression is given by the Bethe formula, which can be heuristically derived as follows. First, we sum the entropic contributions $\Delta S_{\circ+\square \in \circ}$ of each node $\circ$, including the corrections due to its adjacent parity checks $\square \in \circ$. Second, we note that each parity check $\square$ is involved in $k_\square$ terms, with $k_\square$ being the connectivity of $\square$. To count it only once, we therefore subtract $(k_\square - 1)$ times the entropic contribution $\Delta S_\square$ of each parity check $\square$. This leads to

$$s = \frac{1}{N} \left(\sum_{\circ} \Delta S_{\circ+\square \in \circ} - \sum_{\square} (k_\square - 1) \Delta S_\square \right) = \langle \Delta S_{\circ+\square \in \circ} \rangle - \frac{\langle \ell \rangle}{\langle k \rangle} \sum_k c_k (k - 1) \langle \Delta S_\square^{(k)} \rangle, \quad (17)$$

where $\langle \Delta S_{\circ+\square \in \circ} \rangle$ represents the average of $\Delta S_{\circ+\square \in \circ}$ over the nodes $\circ$ and $\langle \Delta S_\square^{(k)} \rangle$ the average of $\Delta S_\square$ over the parity checks $\square$ with connectivity $k_\square = k$; the factor $\langle \ell \rangle / \langle k \rangle$ accounts for the ratio of the number M of parity checks over the number N of nodes.

To compute $\Delta S_{\circ+\square \in \circ}$, we need to know whether the bits of the nodes adjacent to $\circ$ are fixed or not, in the absence of the “cavity node” $\circ$. As the cavity node is connected to its neighbors through parity checks [see Fig. 4(a)], we can decompose the computation in two steps. First, we observe that a given neighboring parity check constrains the value of the cavity node if and only if all the other nodes to which it is connected have themselves their bit fixed *in the absence of the cavity node*. Denoting by ζ the probability of this event and by η the probability for a node to be free *in the absence of one of its adjacent parity check*, we thus have

$$\zeta = \sum_k \frac{kc_k}{\langle k \rangle} [1 - (1 - \eta)^{k-1}] = 1 - \frac{c'(1 - \eta)}{\langle k \rangle}, \quad (18)$$

where $kc_k / \langle k \rangle$ is the probability for a parity check be connected to $k-1$ nodes *in addition to the cavity node* [see Fig. 4(a)] and $1 - (1 - \eta)^{k-1}$ is the probability that at least one of these $k-1$ nodes is free in the absence of the parity check. Next, we observe that the probability for the cavity node to

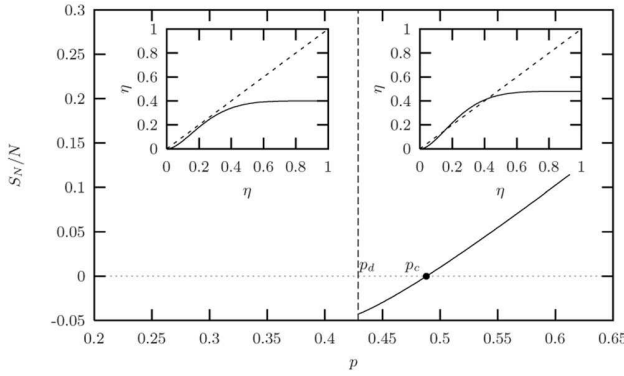


FIG. 5. Reduced entropy vs noise level p for an LDPC code with $k=6$ and $\ell=3$. When $p=0.4 < p_d$ (left inset), $\eta=0$ is the only solution to the cavity equation (24), yielding $s=0$. When $p=0.48 > p_d$ (right inset), two more solutions appear, one of which is stable. The entropy of this solution crosses zero at the critical noise p_c , above which the entropy becomes strictly positive, causing failure of decoding.

be free is the probability that none of its adjacent parity checks is constraining—that is,

$$\nu = p \sum_{\ell} v_{\ell} \zeta^{\ell} = p v(\zeta). \quad (19)$$

In order to close the equations, we also need the probability for the cavity node to be free in the absence of one of its connected parity check [see Fig. 4(c)], which is

$$\eta = p \sum_{\ell} \frac{\ell v_{\ell}}{\langle \ell \rangle} \zeta^{\ell-1} = p \frac{v'(\zeta)}{\langle \ell \rangle}, \quad (20)$$

where $\ell v_{\ell} / \langle \ell \rangle$ represents the probability for a node to be connected to $\ell-1$ parity checks in addition to the one ignored. The “cavity fields” η and ζ , determined by Eqs. (18) and (20), contain all the information needed to evaluate the entropy. Thus $\langle \Delta S_{\circ+\square \in \circ} \rangle$ is given by

$$\langle \Delta S_{\circ+\square \in \circ} \rangle = (\ln 2) [p v(\zeta) - \langle \ell \rangle \zeta]. \quad (21)$$

The first term $(\ln 2) p v(\zeta)$ corresponds to $(\ln 2) \nu$ see [Eq. (19)], the average entropic contribution of a node $\circ$, and the second term $-(\ln 2) \langle \ell \rangle \zeta$ subtracts the entropic reductions of its adjacent parity-check nodes; indeed, they are $\langle \ell \rangle$ on average and each is constraining the cavity node with probability ζ . Similarly, the average entropic reduction due to a parity check alone is

$$\langle \Delta S_{\square}^{(k)} \rangle = -(\ln 2) [1 - (1 - \eta)^k] \quad (22)$$

since $1 - (1 - \eta)^k$ is the probability that at least one of the k connected nodes is free in the absence of the parity check [see Fig. 4(b)]. To sum up, the entropy is determined by the formulas

$$s = (\ln 2) \left[p v \left(1 - \frac{c'(1 - \eta)}{\langle k \rangle} \right) - \frac{\langle \ell \rangle}{\langle k \rangle} [1 - c(1 - \eta) - \eta c'(1 - \eta)] \right], \quad (23)$$

$$\eta = p \frac{v'(1 - c'(1 - \eta) / \langle k \rangle)}{\langle \ell \rangle}. \quad (24)$$

Equation (24) can admit two kinds of solution (see Fig. 5). The first kind, referred to as *ferromagnetic*, describes the situation where decoding is possible, with only one code-word being solution of the decoding CSP: this solution has $\eta=0$ (all bits are fixed to σ^*) and $s=0$. The second kind, referred to as *paramagnetic* (but strictly speaking corresponding to a 1RSB glassy solution), describes the situation where decoding is impossible and has $\eta > 0$. It is found to exist only for p greater than the so-called *dynamical threshold*, denoted by p_d . It is, however, relevant only when associated with a positive entropy, $s > 0$, a condition which defines the *static threshold*, denoted by p_c and satisfying $p_c > p_d$. The static threshold corresponds to the threshold above which decoding is doomed to fail, as confirmed by rigorous studies.

3. Algorithmic interpretation

The cavity method is related to a particular decoding algorithm known as belief propagation (BP). Its principle is the following: starting from a configuration where only the noncorrupted bits are fixed to their values, one goes through each node of the factor graph, checks if its immediate neighboring environment constrains it to a unique value, fixes it to this value if it is the case, and iterates the whole procedure until convergence. At the end, some bits may still not be fixed, which certainly occurs if the decoding CSP has not a unique solution, but if all the bits end up fixed, one is ensured to have correctly decoded. Similar message-passing algorithms can be defined with different channels. They are responsible for the practical interest of LDPC codes as they provide algorithmically efficient decoding (yet suboptimal, as discussed below). With the BEC, these algorithms are particularly easy to analyze thanks to the fact that one can never be fooled by fixing bits to an incorrect value. To perform the analysis of the possible outcomes of the belief propagation algorithm, we can assume without loss of generality that the transmitted message is $(0, \dots, 0)$ (the $\mathbb{Z}_2$ symmetry implies that all codewords are equivalent). We thus start with $\sigma_i = *$ if $i \in \mathcal{E}$ and $\sigma_i = 0$ otherwise. Cavity fields are attributed to each oriented link of the factor graphs and are updated with the following rules, where t indexes iteration steps:

$$h_{i \rightarrow a}^{(t+1)} = \begin{cases} 0 & \text{if } \sigma_i = 0 \text{ or if } u_{b \rightarrow i}^{(t)} = 1 \text{ for some } b \in i - a, \\ * & \text{otherwise,} \end{cases}$$

$$u_{a \rightarrow i}^{(t+1)} = \begin{cases} 1 & \text{if } h_{j \rightarrow a}^{(t)} = 0 \text{ for all } j \in a - i, \\ * & \text{otherwise.} \end{cases} \quad (25)$$

Here, $u_{a \rightarrow i}^{(t)} = 1$ (*) means that the parity check a is constraining (is not constraining) i . $h_{i \rightarrow a}^{(t)} = 0$ (*) means that σ_i is fixed (not determined) to its correct value 0 without taking

into account the constraints due to a . The algorithm is analyzed statistically by introducing

$$\eta^{(t)} = \frac{1}{\langle \ell \rangle N} \sum_{(i,a)} \delta(h_{i \rightarrow a}^{(t)}, 0), \quad \zeta^{(t)} = \frac{1}{\langle k \rangle M} \sum_{(i,a)} \delta(u_{a \rightarrow i}^{(t)}, 1). \quad (26)$$

As suggested by our notations, the evolution for these quantities exactly mimics the derivation of the formulas for the cavity fields, yielding

$$\eta^{(t+1)} = p \frac{v'(\zeta^{(t)})}{\langle \ell \rangle}, \quad \zeta^{(t+1)} = 1 - \frac{c'(1 - \eta^{(t)})}{\langle k \rangle}. \quad (27)$$

The fixed point is given by Eq. (24). When $p < p_d$, the algorithm converges towards the unique, ferromagnetic, fixed point $\eta^{(\infty)} = \zeta^{(\infty)} = 0$ and decoding is successfully achieved. When $p_d < p < p_c$, a paramagnetic fixed point appears in addition to the ferromagnetic fixed point and the iteration leads to this second paramagnetic fixed point. The belief propagation algorithm thus fails to decode above the dynamical threshold p_d , before reaching the static threshold p_c below which no algorithm can possibly be successful (in this sense, BP is suboptimal).

B. Average error exponents

1. Entropic (IRSB) large deviations

The previous section recalled the properties of typical codes subject to typical noise. With finite codewords, $N < \infty$, failure to decode may also be due to atypical noise with unusually destructive effects. This is the purpose of our large deviation approach to investigate such events. We first focus on the simplest case: namely, the computation of the average error exponent where both the codes $\mathcal{C}$ and the noise ξ are treated on the same footing (see Sec. II E). Our procedure to deal with the statistics over atypical factor graphs is an application of the cavity method for large deviations proposed

in [15]. For the sake of simplicity, we restrain ourselves here to regular codes, where nodes and check nodes have both fixed connectivity, ℓ and k , respectively, and defer the generalization to irregular codes to Appendix D.

As explained in Sec. II E, the thermodynamic formalism assigns a Boltzmann weight $e^{xS_N(\mathcal{C}, \xi)}$ to each “configuration” $(\mathcal{C}, \xi)$. The parameter x plays the role of an inverse temperature or, in other words, is a Lagrange multiplier enforcing the value of S_N . Taking the infinite-temperature limit $x=0$ (no constraint on the value of S_N) will thus lead us back to the typical case discussed above.

The cavity equations are as before derived by considering the effect of the addition of a node. As adding a new node, along with its adjacent parity checks, inevitably increases the degrees of the other nodes, strictly restraining to regular graphs is not possible and we must work in a larger framework. Accordingly, we consider ensembles where the degree of parity checks is fixed to k , but where the degree of nodes has a distribution $\{v_L\}$ (meaning that degree L has probability v_L , independently for each node). We will describe the regular ensemble by taking $v_L = \delta_{\ell, L}$ in the final formulas. Adding a new node with ℓ parity checks brings us from an ensemble characterized by v_L to an ensemble characterized by v'_L , with

$$v'_L = \left(1 - \frac{\ell(k-1)}{N}\right)v_L + \frac{\ell(k-1)}{N}v_{L-1} = v_L + \frac{\ell(k-1)}{N}\delta v_L, \quad (28)$$

where $\delta v_L = v_{L-1} - v_L$, since $\ell(k-1)$ nodes have their degree increased by 1. Let denote by $L(s, \{v_L\})$ the rate function for the probability to observe $S_N/N = s$ in an ensemble characterized by $\{v_L\}$ —that is,

$$\mathbb{P}_M[(\mathcal{C}, \xi): S_N(\mathcal{C}, \xi)/N = s | \{v_L\}] \asymp e^{-NL(s, \{v_L\})}. \quad (29)$$

We introduce $P_{\circ+\square\in\circ}^{(\ell)}(\Delta S)$, the probability distribution of the entropy contribution caused by the addition of the new nodes along with its ℓ adjacent parity checks. The passage from N nodes to $N+1$ nodes can then be described by

$$\begin{aligned} \mathbb{P}_{N+1}(s = S/(N+1) | \{v_L\}) &\asymp e^{-(N+1)L(S/(N+1), \{v_L\})} = \sum_{\ell} v_{\ell} \int d\Delta S P_{\circ+\square\in\circ}^{(\ell)}(\Delta S) \mathbb{P}_M[s = (S - \Delta S)/N | \{v_L - \ell(k-1)/N \delta v_L\}] \\ &\asymp \sum_{\ell} v_{\ell} \int d\Delta S P_{\circ+\square\in\circ}^{(\ell)}(\Delta S) e^{-NL[(S - \Delta S)/N, \{v_L - \ell(k-1)/N \delta v_L\}]}. \end{aligned} \quad (30)$$

Expanding for large N , one gets

$$\begin{aligned} \phi_s(x) &= xs - L(s, \{v_L\}) \\ &= \ln \sum_{\ell} v_{\ell} \int d\Delta S P_{\circ+\square\in\circ}^{(\ell)}(\Delta S) e^{x\Delta S + z\ell(k-1)}, \end{aligned} \quad (31)$$

with

$$z = \sum_L \delta v_L \frac{\partial L(s, \{v_L\})}{\partial v_L}. \quad (32)$$

The parameter z is determined by noting that the addition of a new parity check changes the node degree distribution in the same way as in Eq. (28), with $v'_L = v_L + (k/N)\delta v_L$, yielding

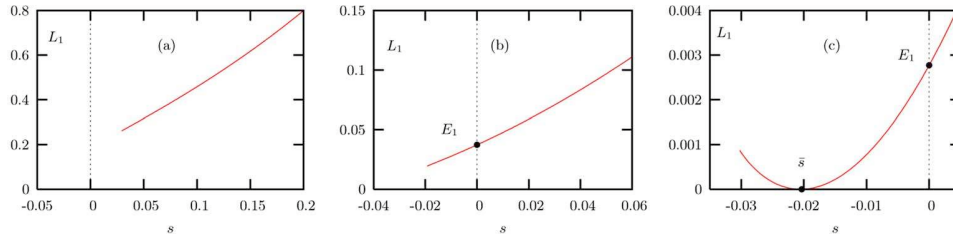


FIG. 6. (Color online) Rate function $L(s)$ as a function of the entropy s , here illustrated with a regular code with $k=6$ and $\ell=3$ (for the BEC channel). The three regimes are represented. (a) $p=0.2 < p_{\text{RSB}}$: the spinodal of the paramagnetic solution is for $s_d > 0$. (b) $p=0.35 \in [p_{\text{RSB}}, p_d]$: the spinodal is now for $s_d < 0$. (c) $p=0.45 \in [p_d, p_c]$: the spinodal is preceded by a minimum (the typical value), with $x_d = \partial_s L(s=s_d) < 0$. The typical dynamical and static transitions can be read on the $s=0$ axis: by definition of p_d and p_c , this equation has a solution $\bar{s}$ for $p > p_d$ and this solution is positive, $\bar{s} > 0$, for $p > p_c$ (not represented here).

$$e^{-NL(S/N, \{v_L\})} \asymp \int d\Delta SP_{\square}(\Delta S) e^{-NL[(S-\Delta S)/N, \{v_L - (k/N)\delta v_L\}]}, \quad (33)$$

where $P_{\square}(\Delta S)$ is the probability of the entropy reduction caused by the addition of a new parity check. Expanding here also for large N leads to an equation for z ,

$$z = -\frac{1}{k} \ln \int d\Delta SP_{\square}(\Delta S) e^{x\Delta S}. \quad (34)$$

Following the same line of reasoning as in the typical case, the two distributions $P_{\square+\square\in\circ}^{(\ell)}$ and $P_{\square}$ can be expressed by means of cavity fields η and ζ . First consider the addition of a node: If the bit of the new node is fixed, either because it was not erased or because one of its adjacent parity checks constrains it, there is an entropic reduction $-\ln 2$ per nonconstraining adjacent parity check and thus a weight 2^{-x} . Otherwise, if the new node is free, which occurs with probability $p\zeta^{\ell}$, the entropy shift is $(\ln 2)(1-\ell)$, giving a weight $2^{x(1-\ell)}$. Taking $v_L = \delta_{L,\ell}$, Eq. (31) therefore reads

$$\phi_s(x) = \ln[(\zeta 2^{-x} + 1 - \zeta)^{\ell} - p(\zeta 2^{-x})^{\ell} + p\zeta^{\ell} 2^{x(1-\ell)}] + \ell(k-1)z, \quad (35)$$

with

$$\zeta = 1 - (1 - \eta)^{k-1}. \quad (36)$$

Similarly, a new parity check removes a degree of freedom if and only if one of its adjacent node is free, which happens with probability $1 - (1 - \eta)^k$, yielding

$$z = -\frac{1}{k} \ln\{1 - [1 - (1 - \eta)^k] + [1 - (1 - \eta)^k]2^{-x}\}. \quad (37)$$

Finally, we obtain a self-consistent equation for η by considering the addition of a new (cavity) node in the absence of one of its adjacent parity checks:

$$\begin{aligned} \eta &= \mathbb{P}(\text{cavity node free}) \\ &\propto \int d\Delta SP_{\circ\rightarrow\square}(\Delta S | \text{cavity node free}) e^{x\Delta S + z(\ell-1)(k-1)}, \end{aligned} \quad (38)$$

$$1 - \eta = \mathbb{P}(\text{cavity node fixed})$$

$$\propto \int d\Delta SP_{\circ\rightarrow\square}(\Delta S | \text{cavity node fixed}) e^{x\Delta S + z(\ell-1)(k-1)}, \quad (39)$$

where $P_{\circ\rightarrow\square}$ corresponds to $P_{\circ+\square\in\circ}^{(\ell-1)}$, taken either under the condition that the cavity node be free or that be is fixed. We obtain

$$\eta = \frac{p2^x(\zeta 2^{-x})^{\ell-1}}{(\zeta 2^{-x} + 1 - \zeta)^{\ell-1} + p(2^x - 1)(\zeta 2^{-x})^{\ell-1}}. \quad (40)$$

Alternatively, these equations can be obtained by differentiation of Eq. (35), which is variational with respect to the cavity η . The large deviation cavity equations (36) and (40) allow us to compute the generating function $\phi_s(x)$ using Eqs. (35) and (37), from which the rate function $L(s | \{v_L = \delta_{L,\ell}\})$ is deduced by Legendre transformation as discussed in Sec. II E.

Again, two kinds of solutions, paramagnetic or ferromagnetic, can be present. For a given value of p , we find that a nontrivial, paramagnetic solution to Eq. (40) exists only for $x \geq x_d(p)$. In agreement with the observation reported in the previous section that the paramagnetic solution typically exists only when $p < p_d$, we have $x_d(p) < 0$ for $p > p_d$ and $x_d(p) > 0$ for $p < p_d$ (the typical case is indeed associated with $x=0$). We obtain the average error exponent by selecting the value of $L(s)$ where $s=0$: our results are illustrated in Fig. 6. By extension of the concept of dynamical threshold p_d , one could define a “dynamical” error exponent as $E_d(p) = L(x_d(p)) = x_d(p)s(x_d(p)) - \phi(x_d(p))$ with $x_d(p)$ corresponding to the temperature of the spinodal for the paramagnetic solution. The relevance of this concept is, however, limited by the fact that the algorithmic interpretation presented in Sec. III A 3 does not extend to large deviations (see also Sec. III C 3).

More interestingly, we find an additional threshold (see Table III), denoted p_{RSB} , below which the equation $s(x)=0$ has no longer a solution (see Fig. 6). This inconsistency of the 1RSB solution is indicative of the presence of a phase transition occurring at some $p_e > p_{\text{RSB}}$. The following section is devoted to computing p_e and describing the nature of the new phase present for $p < p_e$.

TABLE III. Values of some thresholds $p_{1\text{RSB}}$, p_{RS} , p_e , p_d , and p_c for different regular ensembles of LDPC codes on the BEC.

(k, ℓ)	$p_{1\text{RSB}}$	p_{RS}	p_e	p_d	p_c
(4,3)	0.325 262 970 9	0.546 574 881 1	0.606 872 016 6	0.647 425 6494	0.746 009 7025
(6,3)	0.266 856 875 4	0.337 837 464 1	0.349 188 490 2	0.429 439 8144	0.488 150 8842
(6,5)	0.013 008 205 24	0.427 701 036 8	0.714 365 751 3	0.551 003 5344	0.833 315 3204
(10,5)	0.044 128 845 46	0.243 565 689 4	0.334 772 117 6	0.341 550 0230	0.499 490 7179

2. Energetic (RS) large deviations

The previous “entropic (1RSB) approach” attributed errors to the presence of an exponential number of solutions in the decoding CSP. The same assumption was underlying the analysis of the typical case, in Sec. III A 2, where rigorous studies support the conclusions drawn from this hypothesis. This view is also consistent with the phase diagram of XOR-SAT problems to which the encoding CSP belongs. The structure of the well-separated codewords corresponds in this context to a “frozen 1RSB glassy” phase. As p departs from the value $p=1$, however, the decoding CSP deviates increasingly in nature from the initial encoding CSP. As the number of constraints increases (as p decreases), the presence of an exponential number of solutions (glassy phase) in addition to the isolated correct codeword becomes less and less probable. An alternative rare event possibly dominating the probability of error at low p is the presence of a second isolated (ferromagnetic) codeword close to the correct one. This can lead to a new phase transition that has no counterpart in the typical phase diagram, reflected by a nonanalyticity of the error exponent.

In our framework, investigating an alternative source of error requires considering for S_N another quantity than the entropy of the number of solutions. A possible choice, associated with a replica symmetric (RS) ansatz, is the energy E_N of the ground state of the decoding CSP, giving the minimal number of violated parity checks. Ignoring the correct codeword, a second isolated codeword is present if and only if $E_N=0$ (otherwise $E_N>0$). Large deviations of this energy are described by the rate function $L_1(e)$ defined as

$$\mathbb{P}[\xi, C: E_N(\xi, C)/N = e] \asymp e^{-NL_1(e)}. \quad (41)$$

The generating function for the rate function $L_1(e)$, defined by

$$e^{N\phi_e(x)} = \mathbb{E}_{\xi, C}[e^{xE_N(\xi, C)}] = \int de e^{N(xe - L_1(e))}. \quad (42)$$

is given by (see [24] for a similar calculation)

$$\begin{aligned} \phi_1(x) = & \ln \left\{ p \int \prod_{a=1}^{\ell} du_a Q(u_a) \exp \left[-x \left(\sum_{a=1}^{\ell} |u_a| - \left| \sum_{a=1}^{\ell} u_a \right| \right) \right] \right. \\ & + (1-p) \int \prod_{a=1}^{\ell} du_a Q(u_a) \exp \left(-2x \sum_{a=1}^{\ell} \delta_{u_a, -1} \right) \Big\} \\ & - \frac{\ell(k-1)}{k} \ln \int \prod_{i=1}^k dh_i P(h_i) \exp \left[-x \delta \left(\prod_{i=1}^k h_i - 1 \right) \right], \end{aligned} \quad (43)$$

with

$$\begin{aligned} P(h \neq +\infty) \propto & p \int \prod_{a=1}^{\ell-1} du_a Q(u_a) \\ & \times \exp \left[-\frac{x}{2} \left(\sum_{a=1}^{\ell-1} |u_a| - \left| \sum_{a=1}^{\ell-1} u_a \right| \right) \right] \\ & \times \delta \left(h - \sum_{a=1}^{\ell-1} u_a \right), \end{aligned} \quad (44)$$

$$P(h = +\infty) \propto (1-p) \int \prod_{a=1}^{\ell-1} du_a Q(u_a) \exp \left(-x \sum_{a=1}^{\ell-1} \delta_{u_a, -1} \right), \quad (45)$$

$$Q(u) = \int \prod_{i=1}^{k-1} dh_i P(h_i) \delta \left[u - S \left(\prod_{i=1}^{k-1} h_i \right) \right], \quad (46)$$

where $S(x)=1$ if $x>0$, -1 if $x<0$, and 0 if $x=0$. Since u only takes values in $\{-1, 0, +1\}$ and h is restrained to integer values, we can introduce

$$Q(u) = q_+ \delta(u-1) + q_- \delta(u+1) + q_0 \delta(u) \quad (47)$$

and

$$p_+ = \int_{h>0} dh P(h), \quad p_- = \int_{h<0} dh P(h), \quad p_0 = 1 - p_+ - p_-. \quad (48)$$

Our interest is here in zero-energy ground states, described by the limit $x \rightarrow \infty$, where the equations simplify to

$$\phi_e(x=+\infty) = -L(e=0) = \ln[(1-q_-)^\ell + p(1-q_+)^\ell - pq_0^\ell] - \frac{\ell(k-1)}{k} \ln \left[1 - \frac{1}{2}[(p_+ + p_-)^k - (p_+ - p_-)^k] \right], \quad (49)$$

with

$$p_+ \propto (1-q_-)^{\ell-1} - pq_0^{\ell-1}, \quad (50)$$

$$p_- \propto p(1-q_+)^{\ell-1} - pq_0^{\ell-1}, \quad (51)$$

$$p_0 \propto pq_0^{\ell-1}, \quad (52)$$

$$q_+ = \frac{1}{2}[(p_+ + p_-)^{k-1} + (p_+ - p_-)^{k-1}], \quad (53)$$

$$q_- = \frac{1}{2}[(p_+ + p_-)^{k-1} - (p_+ - p_-)^{k-1}], \quad (54)$$

$$q_0 = 1 - (p_+ + p_-)^{k-1}. \quad (55)$$

We find that the only stable solution to these cavity equations satisfies $q_0=p_0=0$, which allows us to further simplify the formulas

$$\phi_e(+\infty) = \ln[q_+^\ell + p(1-q_+)^\ell] - \frac{\ell(k-1)}{k} \ln \left[\frac{1}{2}[1 + (2p_+ - 1)^k] \right], \quad (56)$$

with

$$p_+ = \frac{q_+^{\ell-1}}{q_+^{\ell-1} + p(1-q_+)^{\ell-1}}, \quad (57)$$

$$q_+ = \frac{1}{2}[1 + (2p_+ - 1)^{k-1}]. \quad (58)$$

The resulting RS average error exponent, given by $E_e(p) = -\phi(+\infty)$, is represented in Fig. 7.

We identify the transition p_e as the point where the 1RSB and RS error exponents coincide, which satisfies $p_e > p_{1RSB}$. We find that the RS solution is limited by a spinodal point and is only defined for $p \geq p_{RS}$. While we conjecture that the 1RSB estimate is exact for $p > p_e$, the existence of p_{RS} suggests that either an additional phase transition occurs at some $p'_e > p_{RS}$ or, more radically, that our description of the phase $p < p_e$ is incorrect. The limit case of random codes, however, indicates that the energetic method is valid in the limit $k, \ell \rightarrow \infty$.

3. Limit of random codes

The only limiting case where the average error exponent has been obtained integrally so far is the fully connected limit where $k, \ell \rightarrow \infty$ with $\ell/k = \alpha = 1 - R$ fixed. This limit corresponds to the *random linear model*, where each parity check is connected to each node with probability 1/2. In this limit, the entropic 1RSB approach gives

$$E_s(k, \ell \rightarrow \infty) = L(s=0) = D(1 - R\|p), \quad (59)$$

where $D(q\|p) = q \ln(q/p) + (1-q) \ln[(1-q)/(1-p)]$ is known as the Kullback-Leibler divergence, while the energetic RS approach gives

$$E_e(k, \ell \rightarrow \infty) = -\phi_e(+\infty) = -(R-1) \ln 2 - \ln(1+p) \quad (60)$$

(with $p_+ = 1/1+p$ and $q_+ = 1/2$). The two expressions coincide at the critical noise p_e , with

$$p_e = (1-R)/(1+R). \quad (61)$$

We thus predict the average error exponent of the RLM to be

$$E_1(\text{RLM}) = \begin{cases} (1-R) \ln 2 - \ln(1+p) & \text{if } p < \frac{1-R}{1+R}, \\ D(1-R\|p) & \text{if } \frac{1-R}{1+R} < p < 1-R. \end{cases} \quad (62)$$

This result coincides with the exact expression (see Appendix B for a direct combinatorial derivation), thus validating our approach in this particular case.

As explained above, we are not able to fully account for the small noise regime as soon as k and ℓ are finite, even though the solutions are found to be stable with respect to further replica symmetry breakings in the space of code-words [30]. This does not exclude that a similar replica symmetry breaking occurs in the space of codes. Remarkably, previous attempts reported in the literature have also failed to obtain error exponents in the low p regime.

C. Typical error exponents

1. Cavity equations

The typical error exponent is encoded into a potential $\psi(x, y)$, as defined in Eq. (13). The equations for $\psi(x, y)$ are obtained from the cavity method for large deviations by following very closely the path leading to $\phi(x)$ [31]. As noticed in Sec. II, the formalism with finite y provides a generalization of the average case which is recovered by taking $y=1$, with $\psi(x, y=1) = \phi(x)$. We will therefore only quote our results. In the entropic (1RSB) case, we find

$$\psi_s(x, y) = \ln[(\zeta 2^{-xy} + 1 - \zeta)^\ell - (\zeta 2^{-xy})^\ell + \zeta^\ell (p 2^x + 1 - p)^y 2^{-\ell xy}] - \frac{\ell(k-1)}{k} \ln\{(1-\eta)^k + [1 - (1-\eta)^k] 2^{-xy}\}, \quad (63)$$

with

$$\eta = \frac{\zeta^{\ell-1} (p 2^x)^y 2^{-(\ell-1)xy}}{(\zeta 2^{-xy} + 1 - \zeta)^{\ell-1} - (\zeta 2^{-xy})^{\ell-1} + \zeta^{\ell-1} (p 2^x + 1 - p)^y 2^{-(\ell-1)xy}},$$

$$\zeta = 1 - (1 - \eta)^{k-1}. \quad (64)$$

In the energetic (RS) case with $x = +\infty$, we find

$$\psi_e(x = +\infty, y) = \ln[q_+^\ell + p^y(1 - q_+)^{\ell-1}] - \frac{\ell(k-1)}{k} \ln\left[\frac{1}{2}[1 + (2p_+ - 1)^k]\right], \quad (65)$$

with

$$p_+ = \frac{q_+^{\ell-1}}{q_+^{\ell-1} + p^y(1 - q_+)^{\ell-1}}, \quad (66)$$

$$q_+ = \frac{1}{2}[1 + (2p_+ - 1)^{k-1}]. \quad (67)$$

In each case, from the potential $\psi(x, y)$, the rate function is obtained as $\mathcal{L}(\phi, x) = y\phi - \psi(x, y)$, with $\phi(x) = \partial_y \psi(x, y)$. By definition, a typical code corresponds to a minimum of $\mathcal{L}$, with $\mathcal{L} = 0$, which, when $\mathcal{L}$ is analytical at this minimum, is associated with $y = \partial_\phi \mathcal{L} = 0$.

As a generic feature, we find that $\mathcal{L}(y, x)$ is an increasing function of y for fixed x , going from negative values for $y < y_c(x)$ to positive ones for $y > y_c(x)$. Negative rate functions, as thus obtained, are certainly unphysical. As negative entropies in the usual cavity-replica method, we attribute them to analytical continuations of physical solutions. The simplest way to circumvent them is, as with the frozen 1RSB ansatz in the replica method, to select $y_c(x)$ with $\mathcal{L}(y, x) = 0$. When $y_c(x) < 1$, meaning that $\mathcal{L}(y=1, x) > 0$, we consider that the average exponent is associated with atypical codes and therefore differs from the typical exponent, described by $\mathcal{L}(y_c(x), x) = 0$. Using this criterion, we find that the two exponents indeed differ for the lowest values of p , when $p < p_y$, where $p_y < p_e$ (see Fig. 8 for an illustration). In general the situation is complicated by the fact that the cavity equations may fail to provide solutions in this regime, as already seen in the average case when $p < p_{RS}$ (corresponding here to $y=1$); the random code limit, where this complication is absent, is thus the most instructive.

2. Limit of random codes

In the limit $k, \ell \rightarrow \infty$, we obtain the following results. In the entropic regime, $p > p_e$, the average and typical exponents are found to coincide. This conclusion extends in the energetic regime only for a restricted interval $[p_y, p_e]$. When $p < p_y$, we have $y_c(x) < 1$ and average and typical error ex-

ponents differ. The formula we obtain for the typical error exponent reads

$$E_{\text{typ}}(\text{RLM}) = \begin{cases} -\delta_{GV}(R) \ln p & \text{if } p < p_y, \\ E_{\text{av}}(\text{RLM}) & \text{if } p_y < p < p_c, \end{cases} \quad (68)$$

with

$$p_y = \frac{\delta_{GV}(R)}{1 - \delta_{GV}(R)}. \quad (69)$$

$\delta_{GV}(R)$ denotes the smallest solution to $(R-1)\ln 2 + H(\delta) = 0$, whose interpretation is discussed in Appendix B. This result, which does not seem to have been reported previously in the literature, coincides with the union bound presented in Appendix C, which strongly suggests that it is indeed exact.

For LDPC with finite connectivity, a similar phase diagram is expected. In the entropic regime, we find indeed that average and typical exponents are identical. In the energetic regime, we face the problem that the cavity equations have no solution below some value of p , which precludes us from estimating p_y .

3. Algorithmic implications

The cavity formalism has the attractive property of corresponding formally to message passing algorithms. Based on this analogy, new algorithmic procedures have been systematically proposed to analyze single finite graphs; each time the cavity approach was found to operate at the ensemble level. With a phase transition occurring at the ensemble level, we have, however, here a system where such a correspondence is no longer meaningful. Following the usual procedure, it is indeed straightforward to implement the cavity approach for average error exponent on a single graph, but in the regime $p < p_y$, this algorithm is doomed to fail: for any typical graph, in the limit of large size, the message passing algorithm will yield the average error exponent, which, as we have seen, is distinct for the correct, typical, error exponent.

IV. LDPC CODES OVER THE BSC

A. Definition

We now turn to error exponents for LDPC codes on the binary symmetric channels. One motivation for repeating the analysis with this channel is that it is representative of a broader class of channels, where bits are not simply erased as

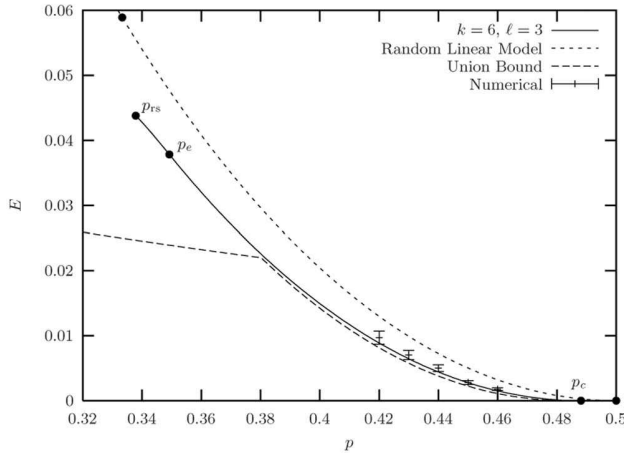


FIG. 7. Average error exponent as a function of the noise level p for the regular code ensemble with $k=6$ and $\ell=3$, on the BEC. Numerical estimates of the error probability, based on 10^6 runs of exact maximum-likelihood decoding (using Gauss elimination) on samples of sizes ranging from $N=500$ to $N=1500$, yield reasonably good estimates of the error exponent using an exponential fit. These numerical results agree well with our theoretical prediction. The union bound (C11) and the random linear limit (62) are also represented for comparison.

with the BEC, but can be *corrupted*, in the sense that their content 0 or 1 is changed to other admissible values. This clearly complicates the decoding as corrupted bits cannot be straightforwardly identified; in fact, with the BSC, no scheme can guarantee to identify the corrupted bits and the receiver is never certain that his decoding is correct. We will, however, see that the overall phase diagram is very similar to that obtained with the BEC.

By definition, maximum-likelihood decoding consists in inferring the most probable realization of the noise *a posteriori*. The *a posteriori* probability can be expressed from the *a priori* probability thanks to Bayes' theorem. If $\mathbf{x}$ denotes the transmitted message and $\mathbf{y}$ the received message, the *a priori* probability to receive $\mathbf{y}$ given $\mathbf{x}$ is

$$Q(\mathbf{y}|\mathbf{x}) = \prod_{i=1}^N (1-p)^{\delta_{x_i y_i}} p^{1-\delta_{x_i y_i}}. \quad (70)$$

To make contact with physical models of disordered systems [12], it is convenient to adopt a spin convention $\sigma_i = (-1)^{x_i}$, $\tau_i = (-1)^{y_i}$, and to rewrite the previous relation as

$$Q(\boldsymbol{\sigma}|\boldsymbol{\tau}) \propto \exp\left(\sum_{i=1}^N h_i \tau_i\right), \quad h_i \equiv h_0 \sigma_i, \quad h_0 \equiv \frac{1}{2} \ln\left(\frac{1-p}{p}\right). \quad (71)$$

This formulation emphasizes the analogy with the random field Ising model [32], a prototypical disordered system. Using the group symmetry of the set of codewords, we can assume, without loss of generality, that the sent codeword is $\boldsymbol{\sigma} = (+1, \dots, +1)$. With this simplification, the random field takes value $h_i = h_0$ with probability $1-p$ and $-h_0$ with prob-

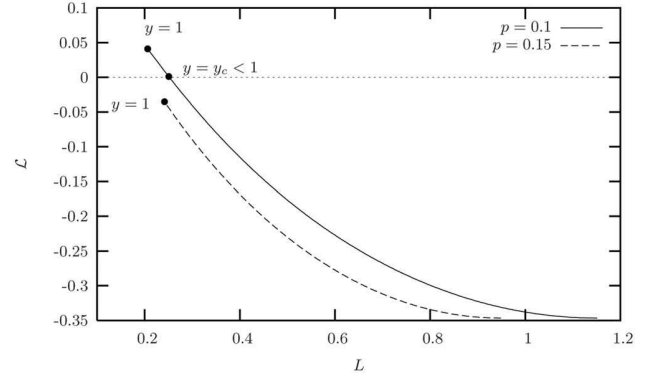


FIG. 8. Rate function $\mathcal{L}(L_e) = \mathcal{L}[-\phi_e(+\infty)]$ of the energetic error exponent for an LDPC code with $k=24$, $\ell=12$ on the BEC. When $p > p_y$ (solid curve), the rate function is negative (and therefore unphysical) for all $0 < y < 1$, entailing that the typical and average error exponents should coincide. When $p < p_y$ (dashed curve), we postulate that the typical error exponent is given by the inverse “freezing temperature” y_c at which the rate function cancels.

ability p . Bayes' formula for the *a posteriori* probability that the message $\boldsymbol{\tau}$ was sent reads

$$\begin{aligned} P(\boldsymbol{\tau}|\boldsymbol{\sigma}) &= \frac{P(\boldsymbol{\sigma}|\boldsymbol{\tau})P(\boldsymbol{\tau})}{\sum_{\boldsymbol{\tau}'} P(\boldsymbol{\sigma}|\boldsymbol{\tau}')P(\boldsymbol{\tau}')} \\ &= \frac{1}{Z(\beta)} \exp\left(\beta \sum_{i=1}^N h_i \tau_i\right) \prod_{a=1}^M \delta(\tau_a = 1), \end{aligned} \quad (72)$$

where τ_a is a shorthand for $\prod_{i \in a} \tau_i$; in the present spin convention, the constraint induced by the parity check a indeed reads $\tau_a = 1$. To continue the analogy with statistical mechanics, we have also introduced a temperature β , called the decoding temperature, whose value is here fixed to $\beta=1$ (Nishimori temperature—see [11]). Given the *a posteriori* probability, the selection of the most probable codeword $\mathbf{d}(\boldsymbol{\sigma})$ can still be done according to different criteria, among which are the following.

(i) Word maximum *a posteriori* (word MAP), where one maximizes the posterior probability in block by taking $\mathbf{d}_{\text{block}}(\boldsymbol{\sigma}) = \arg\max_{\boldsymbol{\tau}} P(\boldsymbol{\tau}|\boldsymbol{\sigma})$. This scheme minimizes the block-error probability $P_{\text{block}} = (1/M) \sum_{\boldsymbol{\tau}} P[\mathbf{d}(\boldsymbol{\sigma}) \neq \boldsymbol{\tau}]$.

(ii) Symbol maximum *a posteriori* (symbol MAP), where one maximizes the posterior probability bit per bit by taking $\mathbf{d}_{\text{bit}}(\boldsymbol{\sigma})_i = \arg\max_{\tau_i} \sum_{\boldsymbol{\tau}_{j \neq i}} P(\boldsymbol{\tau}|\boldsymbol{\sigma})$. This scheme minimizes the bit-error probability $P_{\text{bit}} = (1/M) \sum_{\boldsymbol{\tau}} (1/N) \sum_i P[\mathbf{d}(\boldsymbol{\tau})_i \neq \sigma_i]$.

In physical terms, the word-MAP procedure consists in finding the ground state of the system with partition function $Z(\beta)$ given by the normalization in Eq. (72); this amounts to studying the zero-temperature limit $\beta \rightarrow \infty$. Conversely, symbol MAP is equivalent to taking the sign of the local magnetizations at temperature $\beta=1$,

$$\tau_i^{\text{bit}} = \text{sgn}(\langle \tau_i \rangle) = \text{sgn}\left[\sum_{\boldsymbol{\tau}} \tau_i P(\boldsymbol{\tau}|\boldsymbol{\sigma})\right]. \quad (73)$$

We will treat the two cases in a common framework by considering an arbitrary temperature $\beta \geq 1$.

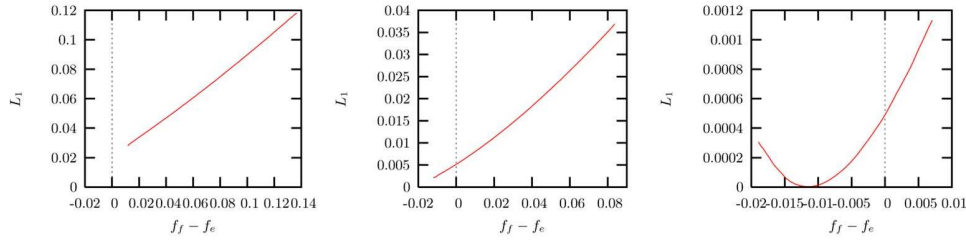


FIG. 9. (Color online) Large deviation rate $L_1(f_f - f_e, s_e = 0)$ as a function of the difference between the ferromagnetic and the nonferromagnetic free energies, here for regular codes with $k=6$ and $\ell=3$ on the BSC. The thresholds are $p_{\text{IRSB}} \approx 0.058$ and $p_c \approx 0.100$. The three regimes are represented. From left to right, $p=0.045$, $p=0.07$, and $p=0.09$.

From the physical perspective, the original codeword is recovered if it dominates the Gibbs measure defined in Eq. (72). This can be expressed by decomposing the partition function $Z(\beta)$ as

$$Z(\beta) = Z_{\text{corr}}(\beta) + Z_{\text{err}}(\beta), \quad Z_{\text{corr}}(\beta) = e^{\beta \sum_i h_i},$$

$$Z_{\text{err}}(\beta) = \sum_{\tau \neq 1} e^{\beta \sum_i h_i \tau_i} \prod_a \delta(\tau_a - 1). \quad (74)$$

We define the corresponding free energies $F_{\text{corr}}(\beta) = -(1/\beta) \ln Z_{\text{corr}}(\beta)$ and $F_{\text{err}}(\beta) = -(1/\beta) \ln Z_{\text{err}}(\beta)$. The first one corresponds physically to a ferromagnetic phase (as with the BEC), while the second will be shown to correspond to either a paramagnetic or a glassy phase, depending on the values of β and p . Decoding is successful if, and only if, the ferromagnetic phase has lower free energy, $F_{\text{corr}} < F_{\text{err}}$. The quantity $S_N(\xi, \mathcal{C})$ introduced in Sec. II E can therefore be defined here as

$$S_N = F_{\text{corr}}(\beta) - F_{\text{err}}(\beta), \quad (75)$$

where the dependence in the noise ξ and the code $\mathcal{C}$ is implicitly understood.

B. Cavity analysis and the 1RSB frozen ansatz

As with the BEC, explicit calculations can be performed by means of the replica or cavity methods. Details can be found in Appendix E, and we only discuss here the points where differences with the BEC arise. For any fixed p , a replica-symmetric calculation, whose derivation follows the derivation of the paramagnetic solution with the BEC, is found to undergo an entropy crisis—i.e., $s_{\text{RS}}(\beta) = \beta^2 \partial_{\beta} f_{\text{RS}}(\beta) < 0$ for $\beta > \beta_g$. This feature is indicative of the presence of a glassy phase and points to the need to break the replica symmetry. The glassy phase of LDPC codes is, however, of the “frozen 1RSB” type, which implies that the glassy free energy f_{err} can be completely inferred from the replica-symmetric solution f_{RS} . This simplicity stems from the “hard” nature of the constraints: changing a bit automatically violates all its surrounding checks, forcing the rearrangement of many variables [33,34]. When the degree of all nodes is $\ell_i \geq 2$, one can indeed show [24] that changing one bit while keeping all checks satisfied requires the rearrangement of an extensive ($\propto N$) number of variables (in the language of [24], factor graphs of LDPC codes have no leaves).

The consequence, expressed in the replica language, is that the 1RSB “states” are reduced to single configurations and thus have zero internal entropy. The 1RSB potential $\phi(\beta, m)$ whose optimization over $m \in [0, 1]$ is predicted to yield f_{err} [20] thus simplifies to $\phi(\beta, m) = f_{\text{RS}}(\beta m)$ [35], since

$$e^{-N\beta m \phi(\beta, m)} \equiv \sum_{\text{states } \alpha} e^{-N\beta m f_{\alpha}(\beta)} = \sum_{\alpha} e^{-N\beta m e_{\alpha}} = e^{-\beta m f_{\text{RS}}(\beta m)}. \quad (76)$$

According to whether one is above or below the freezing temperature β_g^{-1} , defined by

$$s_{\text{RS}}(\beta_g) = \beta_g^2 \partial_{\beta} f_{\text{RS}}(\beta_g) = 0, \quad (77)$$

the free energy $f_{\text{err}}(\beta)$ is given either by $f_{\text{RS}}(\beta)$ (paramagnetic phase) or by $f_{\text{RS}}(\beta_g)$ (glassy phase). This is summarized as follows:

$$f_{\text{err}}(\beta) = \max_{\beta' < \beta} f_{\text{RS}}(\beta') = \begin{cases} f_{\text{RS}}(\beta) & \text{if } \beta < \beta_g, \\ f_{\text{RS}}(\beta_g) & \text{if } \beta > \beta_g. \end{cases} \quad (78)$$

Finally, we note that as in the BEC case, a nonferromagnetic solution $f_{\text{RS}}(\beta)$ exists only for large enough p . The threshold $p_d(\beta)$ giving the smallest noise level at which a nonferromagnetic solution exists is again called the dynamical threshold and can be shown here also to coincide with the dynamical arrest of BP [28].

C. Average error exponent: LDPC codes

In the region relevant for error exponents, where $p < p_c$ and $\beta \geq 1$, the ferromagnetic solution is typically dominant (this is the definition of $p < p_c$) and metastable phases described by f_{err} are typically glassy, since $\beta_g < 1$. Therefore, to compute error exponents, we have to consider $f_{\text{err}}(\beta) = f_{\text{RS}}(\beta_g)$ and not $f_{\text{err}}(\beta) = f_{\text{RS}}(\beta)$. This leads us to introduce an extra temperature β_e distinct from the decoding temperature β , which is to be set to β_g by requiring that the entropy s_{RS} be zero. Similarly, we introduce a ferromagnetic tem-

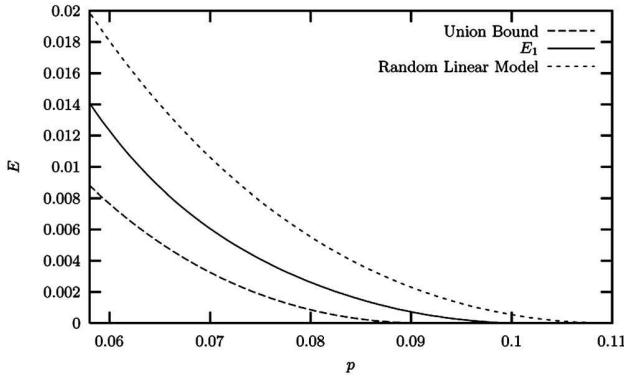


FIG. 10. Average error exponent as a function of the noise level p for the regular code ensemble with $k=6$ and $\ell=3$ through the BSC. Here $p_{\text{IRSB}} \approx 0.058$. The union bound (C17) and the random linear model ($k, \ell \rightarrow \infty$) limit (B14) are also represented for comparison.

perature β_f , set to $\beta_f = \beta$, and define the rate function $L_1(f_e, f_f)$ and its Legendre transform as

$$\begin{aligned} \mathbb{P}[\xi, C: F_{\text{RS}}(\beta_e)/N = f_e, F_{\text{corr}}(\beta_f)/N = f_f] &\asymp e^{-NL_1(f_e, f_f)}, \\ e^{N\phi_1(\beta_e, \beta_f, x_e, x_f)} &= \mathbb{E}_{\xi, C} [e^{-x_e \beta_e F_{\text{RS}}(\beta_e) - x_f \beta_f F_{\text{corr}}(\beta_f)}] \\ &= \int df_e df_f e^{N[-x_e \beta_e f_e - x_f \beta_f f_f - L_1(f_e, f_f)]}. \end{aligned} \quad (79)$$

The potential ϕ_1 contains all the necessary information about both solutions:

$$-\beta_a f_a = \partial_{x_a} \phi_1, \quad s_a = \partial_{x_a} \phi_1 - \frac{\beta_a}{x_a} \partial_{\beta_a} \phi_1, \quad (80)$$

where the index $a=e, f$ corresponds to the two possible phases. For the purpose of computing error exponents, we need only to control $f_e - f_f$ and s_e for all temperatures $\beta_e < \beta$. Note that the ferromagnetic solution f_f has no entropy, $s_f=0$, which is here reflected by the fact that the potential ϕ_1 depends upon β_f and x_f only through $m_f \equiv \beta_f x_f$. These observations allow us to focus on a simplified potential

$$\hat{\phi}(\beta_e, m) = \phi_1\left(\beta_e, x_e = \frac{m}{\beta_e}, m_f = -m\right), \quad (81)$$

which satisfies

$$\partial_m \hat{\phi} = f_f - f_e, \quad \partial_{\beta_e} \hat{\phi} = -ms_e. \quad (82)$$

As with the BEC, the average error exponent is identified with the smallest value of L_1 such that $s_e \geq 0$ and $f_f - f_e \geq 0$. The present formulation is in fact equivalent to the presentation based on the replica method given in [10]. A remarkable consequence of the analysis is that the average error exponent is predicted to be the same for any $\beta \geq 1$. Indeed, both the glassy and the ferromagnetic free energies are temperature independent for $\beta \geq \beta_g$. In particular, symbol and word MAP are predicted to have same error exponents.

Based on the cavity equations given in Appendix E, the potential $\hat{\phi}$ can be computed numerically by population dy-

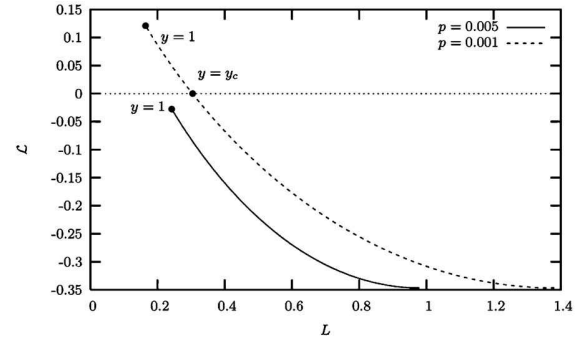


FIG. 11. Rate function $\mathcal{L}(L)$ for the RLM on the BSC with $R=1/2$ and $p=0.005 > p_y$ (solid curve) and $p=0.001 < p_y$ (dashed curve).

namics. As an illustration, we plot in Fig. 9 the rate function $L_1(f_f - f_e, s_e=0)$ for a regular code with $k=6$, $\ell=3$. As in the case of BEC, three regimes can be distinguished, according to the value of p .

- (i) $p < p_{\text{IRSB}}$: no zero-entropy RS solution typically exists and $f_e < f_f$ for the metastable solutions.
- (ii) $p_{\text{IRSB}} < p < p'_d$: no zero-entropy RS solution typically exists but the dominant metastable solutions have $f_e > f_f$.
- (iii) $p'_d < p < p_c$: a zero-entropy RS solution is typically present.

The major difference with the BEC is that the threshold p'_d , defined by $p'_d = p_d(\beta_g(p'_d))$ does not coincide with the dynamical threshold $p_d(\beta)$: indeed here p'_d is defined in relation to the existence of a solution with positive entropy, while, in the framework of BP, the dynamical arrest p_d is related to the existence of a paramagnetic solution at decoding temperature β^{-1} [28]. In Fig. 10, we plot the average error exponent for regular codes with $k=6$, $\ell=3$.

D. Random code limit

1. Average error exponent

As with the BEC, the $k, \ell \rightarrow \infty$ limit can be computed exactly, yielding

$$E_1^{(1)} = L_1(f_f = f_e, s_e = 0) = D(\delta_{GV}(R) \| p), \quad (83)$$

where $\delta_{GV}(R)$ denotes the smallest solution to $R - 1 + H(\delta) = 0$. In this regime, errors are most likely to be caused by large noises driving the received message beyond the typical nearest-codeword distance.

As pointed out in [10], a second ferromagnetic solution is present in this limit (see Appendix E for details), yielding the error exponent

$$E_1^{(2)} = -\ln \frac{1}{2} [1 + 2\sqrt{p(1-p)}] - R \ln 2. \quad (84)$$

Such a solution also exists for finite k, ℓ , but is clearly unphysical (it predicts negative exponents for $k=6$, $\ell=3$). Yet it correctly describes the low p phase (B14) in the $k, \ell \rightarrow \infty$ limit, where failure is caused by the existence of one (or a

few) unusually close codewords. In that sense it plays the same role as the energetic solution in the BEC analysis, with the difference that it is not extensible to any case with finite connectivities. The critical noise p_e below which such a scenario occurs is given by

$$\frac{\sqrt{p_e}}{\sqrt{p_e} + \sqrt{1 - p_e}} = \delta_{GV}(R). \quad (85)$$

We thus predict the average error exponent to be

$$E_1(\text{RLM}) = \begin{cases} D(\delta_{GV}(R) \| p) & \text{if } p < p_e < p_c, \\ -\ln \frac{1}{2}[1 + 2\sqrt{p(1-p)}] - R \ln 2 & \text{if } p < p_e. \end{cases} \quad (86)$$

This expression coincides with the exact result (B14) of the RLM.

2. Typical error exponent

The typical exponent of the RLM can be evaluated using the two-step potential:

$$\psi(\beta_e, m, y) = \mathbb{E}_C[e^{Ny\hat{\phi}(\beta_e, m)}] = \int d\hat{\phi} e^{M[y\hat{\phi} - \mathcal{L}(\hat{\phi}, \beta_e, m)]}. \quad (87)$$

The details of the calculations by the cavity method are given in Appendix E. As in the average case, two distinct solutions appear. The first one is the counterpart of the solution discussed in Sec. IV C. It yields, in the random linear limit,

$$\psi(\beta_e, m, y) = y\hat{\phi}(\beta_e, m). \quad (88)$$

A consequence of the linear dependence on y is that $\hat{\phi}$ always takes the value obtained from the average calculation,

irrespectively of y . Therefore, the average and typical error exponents coincide in this regime and are given by Eq. (83).

This solution is, however, only valid in the high-noise regime ($p > p_e$). As in the average case, for low p , the errors in decoding are dominated by the presence of a subexponential (zero entropy) number of close codewords. The associated solution has for potential

$$\psi(y) = -yL - \mathcal{L} = (R-1)\ln 2 + \ln\{1 + [2\sqrt{p(1-p)}]^y\}. \quad (89)$$

We observe two types of behavior according to the value of p : for $p_y < p < p_e$, $\mathcal{L}(y)$ is negative for $0 \leq y \leq 1$, whereas for $p < p_y$, it crosses 0 at $y_c < 1$ (see Fig. 11). Interpreting, as in the BEC analysis (see Sec. III C 1), negative values of $\mathcal{L}$ as evidence of a glassy transition in the space of codes, we deduce that the typical error exponent is given by $L(y_c)$ when $y_c < 1$, in which case it differs from the average error exponent. To sum up,

$$E_0(\text{RLM}) = \begin{cases} L(y_c) = -\delta_{GV}(R)\ln[2\sqrt{p(1-p)}] & \text{if } p < p_y, \\ L(y=1) = E_1(\text{RLM}) & \text{if } p_y < p < p_c, \end{cases} \quad (90)$$

where the critical noise $p_y(R)$ is a solution of

$$\frac{2\sqrt{p_y(1-p_y)}}{1 + 2\sqrt{p_y(1-p_y)}} = \delta_{GV}(R). \quad (91)$$

This exponent coincides with the RLM limit of the union bound (C18) and is rigorously established [7] to be the correct typical error exponent on the BSC.

V. CONCLUSION

Since Shannon laid the basis for information theory, the analysis of error-correcting codes has been a major subject of study in this field of science [4]. Error-correcting codes aim

at reconstructing signals altered by noise. Their performance is measured by their error probability—i.e., the probability that they fail in accomplishing this task. For block codes, where the messages are taken from a set of 2^M codewords of length N , it is known that when the rate $R=M/N$ is below the channel capacity R_c , the probability of error behaves, in the limit of large N , at best, as $P_e \sim \exp[-NE(R)]$ [4]. This error exponent $E(R)$, also called reliability function, provides a particularly concise characterization of performance.

For a given code ensemble, two classes of error exponents can generally be distinguished, due to the presence of two levels of “disorder,” one associated with the choice of the code itself and a second associated with the realization of the noise. *Average error exponents* correspond to take the error

probability P_e with respect to these two levels simultaneously, while *typical error exponents* refer to fixed, typical, codes.

In the present paper, we tackled the computation of these two error exponents for a particular class of block codes, the low-density parity-check codes, with two particular channels, the binary erasure channel and the binary symmetric channel. We considered decoding under maximum-likelihood decoding, the best conceivable decoding procedure. We framed the problem in terms of large deviations and applied a recently proposed extension of the cavity method designed to probe atypical events in systems defined on random graphs [15]. This method provides an alternative to the replica method used in [10] to address similar problems, with the advantage of being based on explicitly formulated probabilistic assumptions. With respect to this earlier contribution, our work offers several clarifications, notably on the nature of the different phases, and various extensions, notably to the BEC channel. With this particular channel, our results are analytical, and in the high-noise regime, we conjecture them to be exact. Recent mathematical results on the typical phase diagram [36] foster hope for a confirmation of our results in that context.

From a statistical physics perspective, error exponents are interesting for the richness of their phase diagram, which comprises two phase transitions of different natures. These transitions are observed when the level of noise p is varied at fixed rate R (or, equivalently in the special case of random codes, when the rate R is varied at fixed p). Close to the static threshold, for $p_e < p < p_c$, errors are mostly due to the proliferation of many incorrect codewords in the vicinity of the received message. We interpreted this feature in terms of the presence of a glassy phase, and accordingly, we were able to describe this regime by considering a one-step replica symmetry breaking approach. Below p_e , errors become dominated by the effect of single isolated codewords, which we attributed to a transition towards a ferromagnetic state or 1RSB to RS transition. The noise p_e has its counterpart in the “critical rate” R_e of information theory [4], which marks the point below which only bounds on the reliability function are known. The replica-symmetric approach we employed to investigate the regime $p < p_e$ also turns out to be only approximate, except in the limit of infinite connectivity, where we recovered the error exponents of random linear codes [7]. We also described a second transition occurring at $p_y < p_e$, below which atypical codes come to dominate the average exponent, causing it to differ from the typical error exponent. As it takes place in the space of graphs, this is an example of a critical phenomenon whose description is not accessible to the standard cavity method [14], but only to its extension to large deviations [15] (see also [37] for an other example). However, this second transition should be taken with utmost care, as it relies on an approximate ansatz.

The numerous efforts made in the information theory community to account for the low rate regime $R < R_e$ have so far resulted only in upper and lower bounds for the reliability function [6]. Maybe not too surprisingly, this is also the region of the phase diagram where our methods encounter difficulties. Several examples are, however, now available which demonstrate that statistical physics methods can pro-

vide exact solutions to notoriously difficult mathematical problems. The solutions thus obtained generally sharpen our comprehension both of the system at hand and of the techniques themselves, besides often paving the way for rigorous derivations. In the light of some recent such achievements, extending the present statistical physics approach to reach a thorough understanding of error exponents seems to us a valuable challenge.

ACKNOWLEDGMENTS

The work of T.M. was supported in part by the EC through the network MTR 2002-00319 “STIPCO” and the FP6 IST consortium “EVERGROW.” O.R. thanks the Human Frontier Science Program for support.

APPENDIX A: A NOTE ON THE EXPONENTIAL SCALING

The thermodynamic approach is based on the assumption that the leading contribution to the probability of error decays exponentially with N . However, as initially shown by Gallager, for ensembles of LDPC codes, the probability of error decays only polynomially in N to the leading order. In physical terms, this is due to a few codes (whose number is a polynomial in N) which display a second, metastable, ferromagnetic state at a smaller distance from the ground state (corresponding to the correct codeword) than the numerous configurations forming the paramagnetic state.

To overpass this spurious effect in the simplest, yet purely theoretical way, Gallager focused on the so-called “expurgated ensemble” where the half of the codes with smallest minimum distance is disregarded. On this restricted ensemble which excludes the codes with multiple ferromagnetic states, the error probability decays now exponentially in N at the leading order and can be characterized with an average error exponent. Needless to say, this construction only makes sense as a convenient theoretical way to access good codes.

As the large deviation method automatically overlooks any polynomial contribution, its results actually apply to the “expurgated ensemble.” This is, however, only true to the extent that the expurgation does not affect the distribution of graphs in the ensemble (i.e., does not change the distribution of degrees, of loops, etc.). This is presumably the case, as supported by the construction presented in [38], where an expurgated ensemble much tighter than Gallager’s one is defined by explicitly associating to any random code an expurgated code obtained by modifying only a number $O(1)$ of small loops.

APPENDIX B: RANDOM LINEAR MODEL

Definition

A parity-check code is defined by a $M \times N$ matrix A over $\mathbb{Z}_2$ and its codewords are the vectors $\mathbf{x} = (x_1, \dots, x_N)$ satisfying $A\mathbf{x} = 0$. Code ensembles are therefore subsets of the set of

all 2^{MN} possible matrices. Taking this complete set (with all possible matrices having same probability) defines the so-called *random linear model*. In contrast with LDPC codes, since a typical matrix from the RLM is not sparse, the belief propagation algorithm cannot be used to decode. While of little practical interest due to this absence of efficient decoding algorithm, the RLM has, however, two major theoretical advantages, both originating from its “maximally random” nature: typical codes from the RLM saturate the Shannon bounds, and error exponents can be derived rigorously. We review here some of the established results, which we used in the main text as a reference point to compare our nonrigorous results. Error exponents for the RLM are indeed expected to provide upper bounds for error exponents of LDPC ensemble, which are reached only in the limit of infinite connectivity $k, l \rightarrow \infty$ (this limit is similar to that in which p -spin models approach the random energy model when $p \rightarrow \infty$ [27]).

Weight enumerator function

We first characterize the geometry of the space of codewords by means of the so-called *weight enumerator function*. Given a code C with matrix A , this function gives the number $\mathcal{N}_C(d)$ of codewords $\mathbf{x}$ at (Hamming) distance $d=|\mathbf{x}| \equiv \sum_{i=1}^N x_i$ from the origin:

$$\mathcal{N}_C(d) = \sum_{\mathbf{x}} \delta\left(d, \sum_{i=1}^N x_i\right) \delta(A\mathbf{x}, 0), \quad (\text{B1})$$

where the sum is over all codewords and $\delta(x, y)$ enforces the constraint $x=y$. The *average* weight enumerator function is obtained by averaging over the code ensemble and satisfies

$$\bar{\mathcal{N}}(d) \equiv \mathbb{E}_C[\mathcal{N}_C(d)] = \binom{N}{d} 2^{-M} \asymp e^{N\Sigma(R, \delta=d/N)},$$

$$\Sigma(R, \delta) = (R-1)\ln 2 + H(\delta), \quad (\text{B2})$$

where the limit of infinite block length, $N \rightarrow \infty$, is taken with $M=N(1-R)$ and $d=Nx$. The exponent $\Sigma(R, x)$ defines the so-called *average weight enumerator exponent*. A critical distance is the distance $\delta_{GV}(R)$ defined as the smallest $\delta > 0$ such that $\Sigma(R, \delta)=0$. Codewords at distance $d=N\delta$ with $\delta > \delta_{GV}(R)$ proliferate exponentially. On the other hand, the probability of existence of a codeword at distance $d=N\delta$ with $\delta < \delta_{GV}(R)$ is upper-bounded by $\bar{\mathcal{N}}(d)$ and thus decays exponentially with N . Consequently, for any $\epsilon(N)$ such that $\epsilon(N) \rightarrow \infty$ [e.g., $\epsilon(N)=\sqrt{N}$], only an exponentially small fraction of the codes in the ensemble have a minimal nonzero distance $d=N\delta$ smaller than $N\delta_{GV}(R) - \epsilon(N)$. Excluding these “worst” codes from the RLM defines the *expurgated RLM ensemble*.

Average error exponent over the BEC

Due to the group symmetry of the set of codewords, we can assume without loss of generality that the transmitted codeword is $(0, \dots, 0)$. For a given realization of the disorder

due to a BEC, we denote by $E \subset \{1, \dots, N\}$ the subset of erased bits in the received string and d the number of elements in E . If A is the $M \times N$ matrix representing the code, the submatrix $\tilde{A}^E$ induced by A on E defines the decoding CSP problem: decoding is impossible if and only if the kernel of $\tilde{A}^E$ is nonzero. When all matrices A are sampled with uniform probabilities as in the RLM, the submatrices $\tilde{A}^E$ are also represented with uniform probability. Given a noise realization E of magnitude d , the error probability is the probability that a random $M \times d$ matrix $\tilde{A}^E$ is noninjective,

$$\mathbb{E}_C[\mathbb{P}_N^{(B)}(\mathbf{0})] = \sum_{d=0}^N \binom{N}{d} p^d (1-p)^{N-d} \times \mathbb{P}(\exists \mathbf{x} \neq \mathbf{0} \text{ such that } \tilde{A}^E \mathbf{x} = \mathbf{0}). \quad (\text{B3})$$

When $d > M$, $\tilde{A}^E$ is necessarily noninjective. When $d \leq M$, on the other hand, a straightforward inductive argument [8] gives

$$\mathbb{P}(\exists \mathbf{x} \neq \mathbf{0} \text{ such that } \tilde{A}^E \mathbf{x} = \mathbf{0}) = 1 - \prod_{i=0}^{d-1} (1 - 2^{i-M}). \quad (\text{B4})$$

Consequently, the *exact* expression for the average error probability of the RLM reads

$$\mathbb{E}_C[\mathbb{P}_N^{(B)}(\mathbf{0})] = \sum_{d=0}^M \binom{N}{d} p^d (1-p)^{N-d} \left(1 - \prod_{i=0}^{d-1} (1 - 2^{i-M}) \right) + \sum_{d=M+1}^N \binom{N}{d} p^d (1-p)^{N-d}. \quad (\text{B5})$$

In the $N \rightarrow \infty$, this expression can be evaluated by the saddle-point method. When $p < (1-R)/(1+R)$, the dominant contribution comes from the first sum, with

$$\sum_{d=0}^M \binom{N}{d} p^d (1-p)^{N-d} \left(1 - \prod_{i=0}^{d-1} (1 - 2^{i-M}) \right) \asymp e^{-N[(1-R)\ln 2 - \ln(1+p)]} \quad (\text{B6})$$

and typical number of errors $d=N2p/(1+p)$. When $p > (1-R)/(1+R)$ (and $p < 1-R$ to stay below the capacity), the dominant contribution comes from the second sum, with

$$\sum_{d=M+1}^N \binom{N}{d} p^d (1-p)^{N-d} \asymp e^{-ND(1-R\|p)} \quad (\text{B7})$$

and the typical number of errors $d=N(1-R)$. We thus obtain for the average error exponent of the RLM the expression given in Eq. (62),

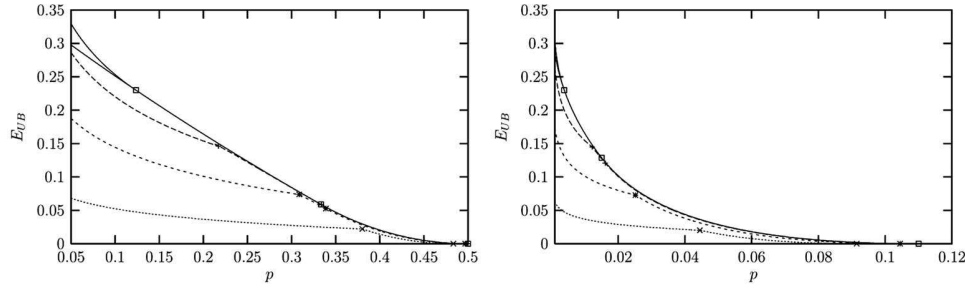


FIG. 12. Expurgated union bounds for the BEC (left) and BSC (right). From bottom to top, $(k, \ell) = (6, 3), (8, 4), (12, 6)$ and the RLM limit, expurgated (top solid curve) and not expurgated (bottom solid curve) with $R = 1/2$. The points indicate the transition between the three regimes, as well as e_{UB} .

$$E_1(\text{RLM}) = \begin{cases} (1-R)\ln 2 - \ln(1+p) & \text{if } p < \frac{1-R}{1+R}, \\ D(1-R\|p) & \text{if } \frac{1-R}{1+R} < p < 1-R. \end{cases} \quad (\text{B8})$$

In physical terms, the transition between the two regimes can be interpreted as a transition between a ferromagnetic (RS) phase and a glassy (1RSB) phase. In the high-noise regime $p > (1-R)/(1+R)$, the error is indeed most probably due to the noise driving the received string into a “glassy phase” of exponentially numerous incorrect codewords, as reflected by the fact that then $\mathbb{P}(\exists \mathbf{x} \neq \mathbf{0} \text{ such that } \tilde{A}^E \mathbf{x} = \mathbf{0}) = 1$. In contrast, in the low-noise regime, $p < (1-R)/(1+R)$, the error is most probably due to the noise driving the received string into a “ferromagnetic phase” where an isolated incorrect codeword happens to be closer than the correct codeword; this is reflected by the fact that $\mathbb{P}(\exists \mathbf{x} \neq \mathbf{0} \text{ such that } \tilde{A}^E \mathbf{x} = \mathbf{0})$ differs from 1 only by an exponentially small term in N , as seen from Eq. (B4).

Average error exponent over the BSC

With the binary symmetric channel, starting again from the transmitted codeword is $(0, \dots, 0)$, the received string $\mathbf{y}$ cannot be decoded if there exists $\mathbf{x} \neq \mathbf{0}$ such that $A\mathbf{x} = \mathbf{0}$ and

$|\mathbf{x} - \mathbf{y}| < |\mathbf{y}|$. Denoting $P_e(\mathbf{y})$ the probability of this event, the probability of error is

$$\mathbb{E}_C[\mathbb{P}_N^{(B)}(\mathbf{0})] = \sum_{d=0}^N \binom{N}{d} p^d (1-p)^{N-d} P_e(\mathbf{y}^{(d)}), \quad (\text{B9})$$

where $\mathbf{y}^{(d)}$ is a generic string of weight d —e.g., $y_i = 1$ if $i \leq d$, $y_i = 0$ if $i > d$. If $d/N > \delta_{GV}(R)$, $P_e(\mathbf{y}^{(d)})$ goes to 1 in the infinite block-length limit. Although no published proof is available in the literature, it is reported as proved [7] that, when $d/N < \delta_{GV}(R)$, $P_e(\mathbf{y}^{(d)})$ is asymptotically equivalent to its union bound approximation (see the following appendix)—i.e.,

$$P_e(\mathbf{y}^{(d)}) \sim \mathbb{E}_C \left[\sum_{\mathbf{x} \neq \mathbf{0}} \theta(d - |\mathbf{x} - \mathbf{y}^{(d)}|) \delta(A\mathbf{x}, \mathbf{0}) \right] \quad (\text{B10})$$

$$\sim \sum_{i=0}^d \mathbb{E}_C[\mathcal{N}_C(i, \mathbf{y}^{(d)})] \quad (\text{B11})$$

$$\sim \mathbb{E}_C[\mathcal{N}_C(d, \mathbf{y}^{(d)})], \quad (\text{B12})$$

where $\mathcal{N}_C(i, \mathbf{y}^{(d)})$ is the number of codewords at distance i from $\mathbf{y}^{(d)}$ and $\theta(x) = 1$ if $x > 0$ and 0 otherwise. Straightforward combinatorics shows that the asymptotic behavior of $\mathbb{E}_C \mathcal{N}_C(i, \mathbf{y}^{(d)})$ is given by the standard weight enumerator exponent $\Sigma(R, i/N)$. In the limit $N \rightarrow \infty$ where $\delta = d/N$ is kept fixed, a saddle-point evaluation leads to the following expression of the average error exponent:

$$E_1(\text{RLM}) = - \max_{\delta < \delta_{GV}} [\Sigma(R, \delta) - D(\delta \| p)] \quad (\text{B13})$$

$$= \begin{cases} (1-R)\ln 2 - \ln[1 + 2\sqrt{p(1-p)}] & \text{if } \frac{\sqrt{p}}{\sqrt{p} + \sqrt{1-p}} < \delta_{GV}(R), \\ D(\delta_{GV}(R) \| p) & \text{otherwise.} \end{cases} \quad (\text{B14})$$

This result with two distinct regime is very similar to that obtained previously for the BEC.

APPENDIX C: UNION BOUNDS

The so-called *union bound exponent* is a rigorous lower bound of the average error exponent in the expurgated ensemble. We show in this appendix how the average weight enumerator exponent of (regular) LDPC codes can be used to derive this union bound exponent, for both the BEC and BSC. We will thus recover results first established by Gallager in [4,39]. In a nutshell, the idea of the union bound is to upper-bound the probability that at least one (bad) codeword causes an error by the sum of the probabilities that each does. Remarkably, this union bound turns out to be tight for the RLM ensemble.

Weight enumerator function

The weight enumerator function [see Eq. (B1) for the definition] of regular LDPC codes with $k=6$ and $\ell=3$ was computed in [4] and reads

$$\begin{aligned} \mathbb{E}_C[\mathcal{N}_C(d)] &= \sum_{\mathbf{x}} \delta(|\mathbf{x}|, d) \mathbb{E}_C[\delta(A\mathbf{x} = \mathbf{0})] \\ &= \binom{N}{d} \mathbb{E}_C[\delta(A\mathbf{x}^{(d)} = \mathbf{0})] \end{aligned} \quad (C1)$$

$$\mathbb{E}_C[\mathcal{N}_C(d = \delta N)] \asymp e^{N\Sigma(k, \ell, \delta)}, \quad (C2)$$

with

$$\Sigma(k, \ell, \delta) = \min_{\mu} \left(2\mu \ell \delta + (1 - \ell) H(\delta) + \frac{\ell}{k} \ln C(\mu) \right), \quad (C3)$$

and

$$C(\mu) = \frac{1}{2} [(1 + e^{-2\mu})^k + (1 - e^{-2\mu})^k]. \quad (C4)$$

We introduce δ_m , the smallest δ such that $\Sigma(k, \ell, \delta) \geq 0$. By construction, the average enumerator exponent in the expurgated ensemble is

$$\Sigma_{\text{exp}}(k, \ell, \delta) = \begin{cases} \Sigma(k, \ell, \delta) & \text{if } \Sigma(k, \ell, \delta) > 0 \text{ (i.e., if } \delta > \delta_m), \\ -\infty & \text{otherwise.} \end{cases} \quad (C5)$$

This expurgated *average* enumerator exponent $\Sigma_{\text{exp}}(k, \ell, \delta)$ is believed to coincide with the *typical* enumerator exponent [40,41].

Union bound for the BEC

Given the set E of erased bits, we want to estimate the probability $P_e(d)$ that the CSP-decoding problem has at least two solutions, when a code C is drawn at random from its ensemble. We call A the matrix characterizing C , $\tilde{A}^E$ the

submatrix induced by A on E , and d the number of erased bits. The union bound consists in the following inequality:

$$P_e(d) = \mathbb{P}(\exists \tilde{\mathbf{x}} \in \{0,1\}^d \neq \mathbf{0} \text{ such that } \tilde{A}^E \tilde{\mathbf{x}} = \mathbf{0}) \quad (C6)$$

$$\leq \min \left[\sum_{\tilde{\mathbf{x}} \neq \mathbf{0}} \mathbb{P}(\tilde{A}^E \tilde{\mathbf{x}} = \mathbf{0}), 1 \right]. \quad (C7)$$

Let $w = |\tilde{\mathbf{x}}|$ and $\mathbf{x}$ be constructed from $\tilde{\mathbf{x}}$ by setting $x_i = \tilde{x}_i$ for $i \in E$, $x_i = 0$ otherwise: $\tilde{\mathbf{x}}$ belongs to the kernel of $\tilde{A}$ if and only if $\mathbf{x}$ belongs to the kernel of A . The probability of the latter event reads

$$\mathbb{E}_C[\mathcal{N}_C(w)] \binom{N}{w}^{-1}. \quad (C8)$$

The error probability is consequently bounded by

$$\mathbb{E}_C[\mathbb{P}_N^{(B)}] = \sum_{d=0}^N \binom{N}{d} p^d (1-p)^{N-d} P_e(d) \quad (C9)$$

$$\begin{aligned} &\leq \sum_{d=0}^N \binom{N}{d} p^d (1-p)^{N-d} \\ &\quad \times \min \left[\sum_{w=0}^d \binom{d}{w} \mathbb{E}_C[\mathcal{N}_C(w)] \binom{N}{w}^{-1}, 1 \right]. \end{aligned} \quad (C10)$$

In the infinite block-length limit, a saddle-point estimate yields, as upper bound for the expurgated average error exponent, the exponent

$$\begin{aligned} E_{\text{exp}}(k, \ell) &\geq E_{UB} \\ &= -\max_{\delta} \left\{ -D(\delta \| p) \right. \\ &\quad \left. + \min \left[\max_{\omega} \left(\Sigma(\omega) + \delta H\left(\frac{\omega}{\delta}\right) - H(\omega) \right), 0 \right] \right\} \\ &= -\max_{\delta < \delta_{UB}} \left\{ -D(\delta \| p) + \max_{\omega > \delta_m} \min_{\mu} \left[\delta H\left(\frac{\omega}{\delta}\right) \right. \right. \\ &\quad \left. \left. + 2\mu \ell \omega - \ell H(\omega) + \frac{\ell}{k} \ln C(\mu) \right] \right\}, \end{aligned} \quad (C11)$$

where $\delta = d/N$, $\omega = w/N$, and δ_{UB} is the largest δ such that $\max_{\omega} (\Sigma(\omega) + \delta H(\frac{\omega}{\delta}) - H(\omega))$ is nonpositive.

As p is varied, three regimes can be distinguished. For small p , the maximum over ω is reached on the boundary δ_m , meaning that errors are dominated by the nearest codewords. For large p instead, the maximum over δ is reached at δ_{UB} , in which case the union bound is simply replaced by 1, physically corresponding to a large number of bad codewords arising from the large amplitude of the noise. Finally, in the intermediate region of p , the extremum is reached in the interior of the (δ, ω) domain. Note that this last regime is not always present when k and ℓ are too small (for $k=6$ and $\ell=3$ in particular). These three regimes are given in the limit $k, \ell \rightarrow \infty$ by

$$E_0(\text{RLM}) = \begin{cases} -\delta_{GV}(R) \ln p & \text{if } p < p_y, \\ (1-R) \ln 2 - \ln(1+p) & \text{if } p_y < p < \frac{1-R}{1+R}, \\ D(1-R \| p) & \text{if } \frac{1-R}{1+R} < p < 1-R, \end{cases} \quad (\text{C12})$$

with p_y defined as in Eq. (69). Union bounds for the BEC are plotted in Fig. 12 for several regular ensembles.

Union bound for the BSC

The union bound for the BSC is derived following the same steps than for the BEC. The counterpart of Eq. (C6) reads

$$P_e(d) = \mathbb{P}(\exists \mathbf{x} \neq 0 \text{ such that } |\mathbf{x} - \mathbf{y}^{(d)}| < d \text{ and } A\mathbf{x} = 0), \quad (\text{C13})$$

where $\mathbf{y}^{(d)}$ is a generic string of weight d . Let $\mathbf{x}$ be a string a weight w and $Q(w, d, g)$ be the probability for $\mathbf{y}^{(d)}$ to be at distance g from $\mathbf{x}$, conditioned on $|\mathbf{y}^{(d)}| = d$:

$$Q(w, d, g) = \binom{w}{(d-g+w)/2} \binom{N-w}{(d+g-w)/2} \binom{N}{d}^{-1}. \quad (\text{C14})$$

The probability for $\mathbf{y}^{(d)}$ to be at distance g from any code-word $\mathbf{x}$ is upper-bounded by

$$\sum_w \mathbb{E}_C[\mathcal{N}_C(w)] Q(w, d, g), \quad (\text{C15})$$

and we can write

$$P_e(d) \leq \min \left[\sum_{w,g} \mathbb{E}_C[\mathcal{N}_C(w)] Q_C(w, d, g), 1 \right] \\ \asymp \min \left[\sum_w \mathbb{E}_C[\mathcal{N}_C(w)] Q_C(w, d, d), 1 \right]. \quad (\text{C16})$$

From this inequality and Eq. (C9), we obtain the union bound for the error exponent via the saddle-point method:

$$E_{\text{exp}}(k, l) \geq E_{UB} = - \max_{\delta} \{ -D(\delta \| p) + \min_{\omega} [\max_{\omega} (\Sigma(\omega) + L(\omega, \delta, \delta)), 0] \} \\ = - \max_{\delta < \delta_{UB}} \left\{ -D(\delta \| p) + \max_{\omega > \delta_m} \min_{\mu} \left[2\mu \ell \omega + (1 - \ell)H(\omega) + \frac{\ell}{k} \ln C(\mu) + L(\omega, \delta, \delta) \right] \right\}, \\ L(\omega, \delta, \gamma) = \omega H\left(\frac{\delta - \gamma + \omega}{2\omega}\right) + (1 - \omega)H\left(\frac{\delta + \gamma - \omega}{2(1 - \omega)}\right) - H(\delta). \quad (\text{C17})$$

As for the BEC, three regimes can be distinguished, according to the value of p . In the limit $k, \ell \rightarrow \infty$, these three regimes are

$$E_0(\text{RLM}) = \begin{cases} -\delta_{GV}(R) \ln[2\sqrt{p(1-p)}] & \text{if } p < p_y, \\ (1-R) \ln 2 - \ln[1 + 2\sqrt{p(1-p)}] & \text{if } p_y < p < p_e, \\ D(\delta_{GV}(R) \| p) & \text{if } p_e < p < \delta_{GV}(R), \end{cases} \quad (\text{C18})$$

where p_y and p_e are given by Eq. (91) and (85)

Union bounds for the BSC are plotted in Fig. 12.

APPENDIX D: IRREGULAR CODES

Definition of the ensemble

In this appendix we discuss the generalization to irregular graphs. We shall only treat the entropic large deviations with the BEC, but our arguments can easily be generalized to the other cases. With irregular codes, it is necessary to specify more precisely the definition of the ensemble. The usual definition is via the degree distributions v_ℓ and c_k . It is, however, possible to define different ensembles having same distribution and sharing the same typical properties, but differing at the level of atypical properties, including error exponents (see also [15] for similar nonequivalences in an other context).

The simplest construction takes all factor graphs with exactly $v_\ell N$ checks of degree ℓ , $c_k M$ variables of degree k , and pick them with uniform probability. Such ensembles are used to build actual codes, and we shall therefore analyze them with some details.

Average error exponent

We revisit the arguments of Sec. III B and emphasize the differences with the regular case.

A crucial modification is the introduction of Lagrange multipliers enforcing the number of nodes of each degree. Call N_ℓ the number of variables of degree ℓ and M_k the number of checks of degree k . Denote $n_\ell = N_\ell/N$ and $m_k = M_k/N$. The rate L_1 is now a function of the n_ℓ and m_k . Its multiple Legendre transform is defined as

$$\phi(x, \{\lambda_\ell\}, \{\nu_k\}) \doteq xs + \sum_{\ell} \lambda_{\ell} n_{\ell} + \sum_k \nu_k m_k - L_1, \quad (\text{D1})$$

with

$$x = \partial_s L_1, \quad \lambda_{\ell} = \partial_{n_{\ell}} L_1, \quad \nu_k = \partial_{m_k} L_1.$$

Let us consider the addition of a new bit. ℓ checks are added along with it, where ℓ is drawn with probability v_ℓ . Each of these checks, in turn, is connected to $k_a - 1$ old bits ($a = 1, \dots, \ell$), where k_a is drawn with probability $k_a c_{k_a} / \langle k \rangle$. Equation (31) is modified in the following way:

$$\begin{aligned} \phi(x, \{\lambda_\ell\}, \{\nu_k\}) = & \ln \sum_\ell v_\ell \sum_{\{k_1, \dots, k_\ell\}} \prod_{a=1}^\ell \frac{k_a c_{k_a}}{\langle k \rangle} \\ & \times \int d\Delta S P_{\square \in \circ}^{(\ell, k_1, \dots, k_\ell)}(\Delta S) \exp \left[x \Delta S + \sum_{a=1}^\ell [(k_a \right. \\ & \left. - 1) z_{k_a} + \nu_{k_a}] + \lambda_\ell \right]. \end{aligned} \quad (\text{D2})$$

The addition of a variable of degree ℓ is reflected by a factor e^{λ_ℓ} and the addition of a check of degree k by a factor e^{μ_k} . Call the k degree the degree of a variable with respect to checks of degree k . Here z_k is related to the increase of k degrees in the ensemble. Let us consider for a moment a more general setting, where the ensemble is determined by the k -degree distributions, denoted by $v_\ell^{(k)}$ [42]. Then z_k is defined by

$$z_k = \sum_\ell \delta v_\ell^{(k)} \frac{\partial L_1(s, \{v_\ell^{(k)}\})}{\partial v_\ell^{(k)}}, \quad (\text{D3})$$

where $\delta v_\ell^{(k)} = v_{\ell-1}^{(k)} - v_\ell^{(k)}$. z_k is obtained in a very similar way as z in Eq. (37):

$$z_k = -\frac{1}{k} \ln \int d\Delta S P_{\square}^{(k)}(\Delta S) e^{x \Delta S + \nu_k}, \quad (\text{D4})$$

where $P_{\square}^{(k)}(\Delta S)$ now depends on the degree k .

The cavity equation (24) is modified in a very similar way as the expression of ϕ_1 in Eq. (D2). The inversion of the Legendre transformation allows one to recover the relevant quantities:

$$s = \partial_x \phi, \quad n_\ell = \partial_{\lambda_\ell} \phi, \quad m_k = \partial_{\nu_k} \phi. \quad (\text{D5})$$

Replacing $P_{\square \in \circ}^{(\ell, k_1, \dots, k_\ell)}(\Delta S)$ and $P_{\square}^{(k)}(\Delta S)$ by their values, we obtain

$$\phi_1 = xs - L_1 = \ln[v(A) + p(2^x - 1)v(B)], \quad (\text{D6})$$

with

$$A = e^{\lambda_\ell} \sum_k \frac{kc_k}{\bar{k}} e^{(k-1)z_k + \nu_k} [2^{-x} + (1 - 2^{-x})(1 - \nu)^k],$$

$$B = 2^{-x} e^{\lambda_\ell} \sum_k \frac{kc_k}{\bar{k}} e^{(k-1)z_k + \nu_k} [1 - (1 - \nu)^{k-1}],$$

$$z_k = -\frac{1}{k} \ln[2^{-x} + (1 - 2^{-x})(1 - \nu)^k] - \frac{\nu_k}{k},$$

$$\nu = \frac{p2^x v'(B)}{v'(A) + p(2^x - 1)v'(B)}.$$

To evaluate L_1 as a function of s , we simply need to tune the parameters λ_ℓ and m_k such that the conditions $n_\ell = v_\ell$ and $m_k = \alpha c_k$ are satisfied.

In Fig. 13, we represent the error exponent for the irregular ensemble with $v(x) = (1/2)x^3 + (1/2)x^4$ and $c(x) = (1/2)x^6 + (1/2)x^8$.

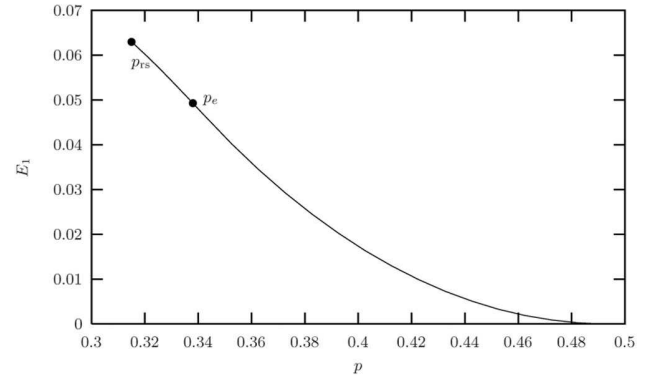


FIG. 13. Average error exponent of a given code as a function of the noise level p for irregular codes with $c_k = (1/2)(\delta_{k,6} + \delta_{k,8})$ and $v_\ell = (1/2)(\delta_{\ell,3} + \delta_{\ell,4})$ through the BEC.

APPENDIX E: CALCULATIONS IN THE BSC

Belief propagation and the Bethe approximation

In this section we write down the BP equations for a given code over the BSC or, equivalently, the cavity equations at the RS level. The expression of the free energy is also given.

The cavity equations read

$$\begin{aligned} p_{\tau_i}^{(i \rightarrow a)} & \propto \prod_{b \in i-a} q_{\tau_i}^{(b \rightarrow i)} e^{-\beta h_i \tau_i}, \\ q_{\tau_i}^{(b \rightarrow i)} & = \sum_{\tau_{b-i}} \prod_{j \in b-i} p_{\tau_j}^{(j \rightarrow b)} \delta[\tau_b = 1]. \end{aligned} \quad (\text{E1})$$

$p_{\tau_i}^{(i \rightarrow a)}$ is the probability that the variable i takes the value τ_i in the absence of a , and $q_{\tau_i}^{(b \rightarrow i)}$ is proportional to the probability that the variable i takes the value τ_i when connected to b only.

Denoting $p_{\tau_i}^{(i \rightarrow a)} = e^{\beta h_{i \rightarrow a} \tau_i} / \cosh \beta h_{i \rightarrow a}$ and $q_{\tau_i}^{(b \rightarrow i)} = e^{\beta u_{b \rightarrow i} \tau_i} / \cosh \beta u_{b \rightarrow i}$, the cavity equations simplify to

$$h_{i \rightarrow a} = \hat{h}(h_i, \{u_{b \rightarrow i}\}) \equiv h_i + \sum_{b \in i-a} u_{b \rightarrow i},$$

$$u_{b \rightarrow i} = \hat{u}(\{h_{j \rightarrow b}\}) \equiv \frac{1}{\beta} \operatorname{arctanh} \left(\prod_{j \in b-i} \tanh \beta h_{j \rightarrow b} \right). \quad (\text{E2})$$

The local magnetization is given by $\langle \sigma_i \rangle = \tanh \beta H_i$, with $H_i = h_i + \sum_{a \in i} u_{a \rightarrow i}$. The Bethe approximation to the free energy reads

$$F_{\text{RS}}(\beta) = \sum_i \Delta F_i - \sum_a (k_a - 1) \Delta F_a, \quad (\text{E3})$$

with

$$\begin{aligned} \Delta F_i & = \Delta F_{\square \in \circ}(\{u_{a \rightarrow i}\}) \equiv \frac{1}{\beta} \sum_{a \in i} \ln[2 \cosh(\beta u_{a \rightarrow i})] \\ & - \frac{1}{\beta} \ln \left[2 \cosh \left(\beta h_i + \beta \sum_{a \in i} u_{a \rightarrow i} \right) \right], \end{aligned}$$

$$\Delta F_a = \Delta F_{\square}(\{h_{i \rightarrow a}\}) \equiv -\frac{1}{\beta} \ln \left(\frac{1 + \prod_{i \in a} \tanh \beta h_{i \rightarrow a}}{2} \right). \quad (\text{E4})$$

Define

$$P(h) = \frac{1}{N\langle \ell \rangle} \mathbb{E}_C \left[\sum_{(i,a)} \delta(h - h_{i \rightarrow a}) \right],$$

$$Q(u) = \frac{1}{N\langle \ell \rangle} \mathbb{E}_C \left[\sum_{(i,a)} \delta(u - u_{a \rightarrow i}) \right]. \quad (\text{E5})$$

Averaging (E1) over the codes, the noise, and the edges, we obtain the self-consistency equations

$$P(h) = \sum_{\ell} \frac{\ell v_{\ell}}{\langle \ell \rangle} \int \prod_{a=1}^{\ell-1} du_a Q(u_a) \langle \delta[h - \hat{h}(h_{\xi}, \{u_a\})] \rangle_{h_{\xi}}, \quad (\text{E6})$$

$$Q(u) = \sum_k \frac{kc_k}{\langle k \rangle} \int \prod_{i=1}^{k-1} P(h_i) \delta[u - \hat{u}(\{h_i\})], \quad (\text{E7})$$

where $h_{\xi} = h_0$ with probability $1-p$ and $-h_0$ with probability p . The RS free energy reads

$$f_{RS}(\beta) = \sum_{\ell} v_{\ell} \int \prod_{a=1}^{\ell} du_a Q(u_a) \langle \Delta F_{\square+\square \in \circ} (h_{\xi}, \{u_a\}) \rangle_{h_{\xi}} - \sum_k c_k (k-1) \int \prod_{i=1}^k dh_i P(h_i) \Delta F_{\square}(\{h_i\}). \quad (\text{E8})$$

Large deviations

As in the BEC, we study the statistics of BP over the codes, under the measure $\propto \exp[-x_f \beta_f F_{\text{corr}}(\beta_f) - x_e \beta_e F_{\text{RS}}(\beta_e)]$. The large deviation cavity equations read, for a regular code,

$$P(h) \propto \int \prod_{a=1}^{\ell-1} du_a Q(u_a) \frac{\left\langle \delta\left(h - h_{\xi} - \sum_{a=1}^{\ell-1} u_a\right) e^{\beta_f x_f h_{\xi}} \left\{ 2 \cosh\left[\beta_e \left(h_{\xi} + \sum_{a=1}^{\ell-1} u_a\right)\right] \right\}^{x_e} \right\rangle_{h_{\xi}}}{\prod_{a=1}^{\ell-1} [2 \cosh(\beta_e u_a)]^{x_e}},$$

$$Q(u) = \int \prod_{i=1}^{k-1} dh_i P(h_i) \delta\left[u - \frac{1}{\beta} \operatorname{arctanh}\left(\prod_{i=1}^{k-1} \tanh(\beta_p h_i)\right)\right], \quad (\text{E9})$$

and the potential

$$\phi(\beta_f, \beta_e, x_f, x_e) = \ln \int \prod_{a=1}^{\ell} du_a Q(u_a) \frac{\left\langle e^{\beta_f x_f h_{\xi}} \left\{ 2 \cosh\left[\beta_e \left(h_{\xi} + \sum_{a=1}^{\ell} u_a\right)\right] \right\}^{x_e} \right\rangle_{h_{\xi}}}{\prod_{a=1}^{\ell} [2 \cosh(\beta_e u_a)]^{x_e}} - \frac{\ell}{k} (k-1) \ln \int \prod_{i=1}^k dh_i P(h_i) \left[\frac{1 + \prod_{i=1}^k \tanh(\beta_e h_i)}{2} \right]^{x_e}. \quad (\text{E10})$$

The solution to (E9) is obtained numerically. In the limit $k, \ell \rightarrow \infty$, this solution simplifies:

$$Q(u) = \delta(u), \quad P(h) = (1-p) \delta(h - h_0) + p \delta(h + h_0), \quad (\text{E11})$$

yielding the error exponent (83).

Another solution, called “type I” in [10], also exists:

$$Q(u) = \eta \delta_{+\infty}(u) + (1-\eta) \delta_{-\infty}(u),$$

$$P(h) = \nu \delta_{+\infty}(h) + (1-\nu) \delta_{-\infty}(h), \quad (\text{E12})$$

with

$$\nu = \frac{\eta^{\ell-1}}{\eta^{\ell-1} + (1-\eta)^{\ell-1} \langle e^{-2y h_0 \sigma} \rangle_{\sigma}}, \quad \eta = \frac{1}{2} [1 + (2\nu - 1)^{k-1}]. \quad (\text{E13})$$

We automatically have $s_p=0$, and the condition $f_p=f_f$ implies $m=\beta_e x_e=1/2$. Then the rate function reads

$$L_1(f_p=f_f) = -\phi = -\ln[\eta^{\ell} + (1-\eta)^{\ell} \langle e^{-h_0 \sigma} \rangle_{\sigma}] - \frac{\ell}{k} (k-1) \ln \left[\frac{1}{2} [1 + (2\nu - 1)^k] \right]. \quad (\text{E14})$$

This solution (E12) is numerically unstable, and the rate function thus obtained is clearly unphysical. However, for $k, \ell \rightarrow \infty$, $\ell/k=1-R$, we have $\eta=\nu=1/2$ and the resulting rate function

$$L_1(f_p=f_f) = -\ln \frac{1}{2} [1 + 2\sqrt{p(1-p)}] - R \ln 2 = \ln 2 [R_0(p) - R] \quad (\text{E15})$$

coincides with the error exponent of the RLM in the low- p regime (B14).

Two-step large deviations

The potential $\psi(\beta_e, m, y)$ defined in Eq. (87) is obtained by extremizing the following expression with respect to $P(h)$ and $Q(u)$:

$$\begin{aligned} \psi(\beta_e, m, y) = & \ln \int \prod_{a=1}^{\ell} du_a Q(u_a) \\ & \times \left\langle \frac{e^{-m h_{\xi}} \left\{ 2 \cosh \left[\beta_e (h_{\xi} + \sum_{a=1}^{\ell} u_a) \right] \right\}^{m/\beta_e}}{\prod_{a=1}^{\ell} [2 \cosh(\beta_e u_a)]^{m/\beta_e}} \right\rangle_{h_{\xi}}^y \\ & - \frac{\ell}{k} (k-1) \ln \int \prod_{i=1}^k dh_i P(h_i) \\ & \times \left[\frac{1 + \prod_{i=1}^k \tanh(\beta_e h_i)}{2} \right]^{ym/\beta_e}. \end{aligned} \quad (\text{E16})$$

We can only handle this calculation in the $k, \ell \rightarrow \infty$ limit. Equations (E11) are still a solution in this case and yield

$$\psi(\beta_e, m, y) = y \hat{\phi}(\beta_e, m), \quad (\text{E17})$$

where $\hat{\phi}(\beta_e, m)$ is obtained from the average case. Therefore, the typical exponent is the same as the average error exponent in the high- p regime.

There also exists a counterpart of solution (E12), which gives

$$\begin{aligned} \psi(\beta_e, m, y) = & (R-1) \ln 2 + \ln \{ 1 + [(1-p)^{1-m} p^m \\ & + p^{1-m} (1-p)^m]^y \}. \end{aligned} \quad (\text{E18})$$

The condition $\partial_m \psi = 0$ is again enforced by setting $m = 1/2$. Thus we get

$$\psi(y) = -yL - \mathcal{L} = (R-1) \ln 2 + \ln \{ 1 + [2\sqrt{p(1-p)}]^y \}. \quad (\text{E19})$$

This expression yields the rate function $\mathcal{L}(L)$ by inverse Legendre transformation.

-
- [1] C. E. Shannon, Bell Syst. Tech. J. **27**, 379 (1948); **27**, 623 (1948).
 - [2] C. Berrou, A. Glavieux, and P. Thitimajshima, in *Proceedings of the IEEE International Conference on Communications, Geneva, 1993* (IEEE, New York, 1993), pp. 1064–1070.
 - [3] D. J. C. MacKay, *Information Theory, Inference, and Learning Algorithms* (Cambridge University Press, Cambridge, England, 2003).
 - [4] R. G. Gallager, IRE Trans. Inf. Theory **IT-8**, 21 (1962).
 - [5] S. Verdú, IEEE Trans. Inf. Theory **44**, 2057 (1998).
 - [6] E. R. Berlekamp, Not. Am. Math. Soc. **49**, 17 (2002).
 - [7] A. Barg and G. D. Forney, Jr., IEEE Trans. Inf. Theory **48**, 2568 (2002).
 - [8] C. Di, D. Proietti, I. E. Telatar, R. L. Urbanke, and T. J. Richardson, IEEE Trans. Inf. Theory **48**, 1570 (2002).
 - [9] A. Amraoui, A. Montanari, T. Richardson, and R. Urbanke, e-print cs.IT/0406060.
 - [10] N. S. Skantzos, J. van Mourik, D. Saad, and Y. Kabashima, J. Phys. A **36**, 11131 (2003).
 - [11] H. Nishimori, *Statistical Physics of Spin Glasses and Information Processing: An Introduction* (Oxford University Press, Oxford, 2001).
 - [12] N. Sourlas, Nature (London) **339**, 693 (1989).
 - [13] F. den Hollander, *Large Deviations, Fields Institute Monographs No. 14* (American Mathematical Society, Providence, RI, 2000).
 - [14] M. Mézard and G. Parisi, Eur. Phys. J. B **20**, 217 (2001).
 - [15] O. Rivoire, J. Stat. Mech.: Theory Exp. 2005, P07004.
 - [16] T. Mora and O. Rivoire, e-print cs.IT/0605130.
 - [17] T. M. Cover and J. A. Thomas, *Elements of Information Theory* (Wiley, New York, 1991).
 - [18] R. M. Tanner, IEEE Trans. Inf. Theory **27**, 533 (1981).
 - [19] B. Bollobás, *Random Graphs*, 2nd ed. (Cambridge University Press, Cambridge, England, 2001).
 - [20] M. Mézard, G. Parisi, and M. A. Virasoro, *Spin-Glass Theory and Beyond, Vol. 9 of Lecture Notes in Physics* (World Scientific, Singapore, 1987).
 - [21] C. H. Papadimitriou and K. Steiglitz, *Combinatorial Optimization, Algorithms and Complexity* (Prentice-Hall, Englewood Cliffs, NJ, 1982).
 - [22] F. Ricci-Tersenghi, M. Weigt, and R. Zecchina, Phys. Rev. E **63**, 026702 (2001).
 - [23] S. Cocco, O. Dubois, J. Mandler, and R. Monasson, Phys. Rev. Lett. **90**, 047205 (2003).
 - [24] M. Mézard, F. Ricci-Tersenghi, and R. Zecchina, J. Stat. Phys. **111**, 505 (2003).
 - [25] M. Mézard and G. Parisi, J. Stat. Phys. **111**, 1 (2003).
 - [26] Y. Kabashima and D. Saad, J. Phys. A **37**, R1 (2004).
 - [27] A. Montanari, Eur. Phys. J. B **23**, 121 (2001).
 - [28] S. Franz, M. Leone, A. Montanari, and F. Ricci-Tersenghi, Phys. Rev. E **66**, 046120 (2002).
 - [29] M. Mézard, M. Palassini, and O. Rivoire, Phys. Rev. Lett. **95**, 200202 (2005).
 - [30] A. Montanari and F. Ricci-Tersenghi, Eur. Phys. J. B **33**, 339 (2003).
 - [31] Contrary to what indicates the last equations of [15], the nature of the order parameter is unchanged when additional levels of disorder are taken into account. The reason is that the cavity method encodes in a unique spatial distribution both the statistics over the nodes of a single graph and the statistics over the graphs in an ensemble. The discrimination between the two levels is done only through the unequal weighting attributed to the different nodes, as controlled by the two independent temperatures x and y .
 - [32] T. Nattermann, in edited by A. P. Young *Spin Glasses and Random Fields* (World Scientific, Singapore, 1998).

- [33] A. Montanari and G. Semerjian, Phys. Rev. Lett. **94**, 247201 (2005).
- [34] A. Montanari and G. Semerjian, J. Stat. Phys. **124**, 103 (2006).
- [35] O. C. Martin, M. Mézard, and O. Rivoire, J. Stat. Phys. **P09006**, 2005.
- [36] C. Measson, A. Montanari, T. Richardson, and R. Urbanke, e-print cs.IT/0410028.
- [37] O. Rivoire and J. Barré, Phys. Rev. Lett. **97**, 148701 (2006).
- [38] J. van Mourik and Y. Kabashima, e-print cond-mat/0310177.
- [39] R. G. Gallager, *Information Theory and Reliable Communication* (Wiley, New York, 1968).
- [40] S. Condamin, <http://www.inference.phy.cam.ac.uk/condamin/report.ps>
- [41] C. Di, A. Montanari, and R. Urbanke, in *Proceedings of the International Symposium on Information Theory, 2004* (IEEE, New York, 2004), p. 102.
- [42] In our case $v_{\ell}^{(k)} = \sum_{\ell' \geq \ell} v_{\ell'} \binom{\ell'}{\ell} c_k^{\ell} (1 - c_k)^{\ell' - \ell}$.

Bibliographie

- [ABM01] Dimitris Achlioptas, Paul Beame, and Michael S. O. Molloy. A sharp threshold in proof complexity. In *ACM Symposium on Theory of Computing*, pages 337–346, 2001.
- [ACIM01] Dimitris Achlioptas, Arthur Chtcherba, Gabriel Istrate, and Christopher Moore. The phase transition in 1-in- k sat and nae 3-sat. In *SODA '01 : Proceedings of the twelfth annual ACM-SIAM symposium on Discrete algorithms*, pages 721–722, Philadelphia, PA, USA, 2001. Society for Industrial and Applied Mathematics.
- [AF99] Dimitris Achlioptas and Ehud Friedgut. A sharp threshold for k -colorability. *Random Struct. Algorithms*, 14(1) :63–70, 1999.
- [AHS87] David H. Ackley, Geoffrey E. Hinton, and Terrence J. Sejnowski. A learning algorithm for boltzmann machines. pages 522–533, 1987.
- [Ald01] D. J. Aldous. The $\zeta(2)$ limit in the random assignment problem. *Rand. Struct. Algo.*, 18 :381–418, 2001.
- [AM02] Dimitris Achlioptas and Cristopher Moore. The asymptotic order of the random k -sat threshold. *focs*, 00 :779, 2002.
- [ANP05] D. Achlioptas, A. Naor, and Y. Peres. Rigorous location of phase transitions in hard optimization problems. *Nature*, 435 :759–764, 2005.
- [AP04] D. Achlioptas and Y. Peres. The threshold for random k -sat is $2^k \log 2 - O(k)$. *Journal of the AMS*, 17 :947–973, 2004.
- [ART06] Dimitris Achlioptas and Federico Ricci-Tersenghi. On the solution-space geometry of random constraint satisfaction problems. In *STOC '06 : Proceedings of the thirty-eighth annual ACM symposium on Theory of computing*, pages 130–139, New York, NY, USA, 2006. ACM Press.
- [Bal83] Roger Balian. *Du Microscopique au Macroscopique, Cours de physique statistique de l'Ecole polytechnique*. Ellipses, 1983.
- [BB04] J. P. Bouchaud and G. Biroli. On the Adam-Gibbs-Kirkpatrick-Thirumalai-Wolynes scenario for the viscosity increase of glasses. *J. Chem. Phys.*, 121 :7347–7354, 2004.

- [BBCZ05] D. Battaglia, A. Braunstein, J. Chavas, and R. Zecchina. Source coding by efficient selection of ground-state clusters. *Phys. Rev. E*, 72(1) :015103, July 2005.
- [BBDR05] Julien Barré, Freddy Bouchet, Thierry Dauxois, and Stefano Ruffo. Large deviation techniques applied to systems with long-range interactions. *J. Stat. Phys.*, 119 :677, 2005.
- [BBL05] J. Barré, A. R. Bishop, T. Lookman, and A. Saxena. On adaptability and intermediate phase in randomly connected networks. *Phys. Rev. Lett.*, 94 :208701, 2005.
- [Ber02] E. R. Berlekamp. The performance of block codes. *Notices of the AMS*, pages 17–22, January 2002.
- [Bet35] H. A. Bethe. Statistical physics of superlattices. *Proc. Roy. Soc. London A*, 150 :552–575, 1935.
- [BGT93] C. Berrou, A. Glavieux, and P. Thitimajshima. Near shannon limit error-correcting coding : Turbo codes. In *Proc. IEEE International Conference on Communications*, pages 1064–1070, 1993.
- [BJ02] A. Barg and G. D. Forney Jr. Random codes : minimum distances and error exponents. *IEEE Trans. Inform. Theory*, 48 :2568–2573, 2002.
- [BM04] David Burshtein and Gadi Miller. Asymptotic enumeration methods for analyzing ldpc codes. *IEEE Transactions on Information Theory*, 50(6) :1115–1131, 2004.
- [BMW00] G. Biroli, R. Monasson, and M. Weigt. A variational description of the ground state structure in random satisfiability problems. *Eur. Phys. J. B*, 14 :551, 2000.
- [BMWZ02] A. Braunstein, M. Mezard, M. Weigt, and R. Zecchina. Constraint Satisfaction by Survey Propagation. *ArXiv Condensed Matter e-prints*, December 2002.
- [BMZ05] A. Braunstein, M. Mézard, and R. Zecchina. Survey propagation : An algorithm for satisfiability. *Random Struct. Algorithms*, 27(2) :201–226, 2005.
- [Bol01] B. Bollobás. *Random graphs*. Cambridge University Press, second edition, 2001.
- [BZ04] A. Braunstein and R. Zecchina. Survey propagation as local equilibrium equations. *Journal of Statistical Mechanics : Theory and Experiment*, 6 :P06007, June 2004.
- [BZ06] A. Braunstein and R. Zecchina. Learning by Message Passing in Networks of Discrete Synapses. *Physical Review Letters*, 96(3) :030201, January 2006.

- [CC06] M. Chertkov and V. Y. Chernyak. Loop series for discrete statistical models on graphs. *Journal of Statistical Mechanics : Theory and Experiment*, 6 :P06009, 2006.
- [CDMM03] S. Cocco, O. Dubois, J. Mandler, and R. Monasson. Rigorous decimation-based construction of ground pure states for spin glass models on random lattices. *Phys. Rev. Lett.*, 90 :047205, 2003.
- [CF86] Ming-Te Chao and John Franco. Probabilistic analysis of two heuristics for the 3-satisfiability problem. *SIAM J. Comput.*, 15(4) :1106–1118, 1986.
- [CF90] Ming-Te Chao and John Franco. Probabilistic analysis of a generalization of the unit-clause literal selection heuristics for the k satisfiability problem. *Inf. Sci.*, 51(3) :289–314, 1990.
- [CKT91] Peter Cheeseman, Bob Kanefsky, and William M. Taylor. Where the Really Hard Problems Are. In *Proceedings of the Twelfth International Joint Conference on Artificial Intelligence, IJCAI-91, Sidney, Australia*, pages 331–337, 1991.
- [CLP⁺06a] L. Correale, M. Leone, A. Pagnani, M. Weigt, and R. Zecchina. Core Percolation and Onset of Complexity in Boolean Networks. *Physical Review Letters*, 96(1) :018101, January 2006.
- [CLP⁺06b] L. Correale, M. Leone, A. Pagnani, M. Weigt, and R. Zecchina. The computational core and fixed point organization in Boolean networks. *Journal of Statistical Mechanics : Theory and Experiment*, 3 :P03002, March 2006.
- [CM05] S. Ciliberti and M. Mézard. The theoretical capacity of the Parity Source Coder. *Journal of Statistical Mechanics : Theory and Experiment*, 10 :P10003, October 2005.
- [CMMS04] S. Cocco, R. Monasson, A. Montanari, and G. Semerjian. Approximate analysis of search algorithms with “physical” methods. In A. Percus G. Istrate, C. Moore, editor, *Phase transitions and Algorithmic complexity*. 2004.
- [CMZ05a] S. Ciliberti, M. Mézard, and R. Zecchina. Lossy Data Compression with Random Gates. *Physical Review Letters*, 95(3) :038701, July 2005.
- [CMZ05b] S. Ciliberti, M. Mezard, and R. Zecchina. Message passing algorithms for non-linear nodes and data compression. *ArXiv Condensed Matter e-prints*, August 2005.
- [CNRTZ03] Tommaso Castellani, Vincenzo Napolano, Federico Ricci-Tersenghi, and Riccardo Zecchina. Bicoloring random hypergraphs. *J.PHYS.A*, 36 :11037, 2003.

- [Con02] S. Condamin. Study of the weight enumerator function for a gallager code. 2002. <http://www.inference.phy.cam.ac.uk/condamin/report.ps>.
- [Coo71] Stephen A. Cook. The complexity of theorem-proving procedures. In *STOC '71 : Proceedings of the third annual ACM symposium on Theory of computing*, pages 151–158, New York, NY, USA, 1971. ACM Press.
- [CS88] Vašek Chvátal and Endre Szemerédi. Many hard examples for resolution. *J. ACM*, 35(4) :759–768, 1988.
- [CT91] T. M. Cover and J. A. Thomas. *Elements of information theory*. Wiley, New-York, 1991.
- [DB97] O. Dubois and Y. Boufkhad. A general upper bound for the satisfiability threshold of random r-sat formulae. *J. Algorithms*, 24(2) :395–420, 1997.
- [Der80] B. Derrida. Random-energy model : Limit of a family of disordered models. *Phys. Rev. Lett*, 45 :79–82, 1980.
- [Der81] B. Derrida. Random-energy model : An exactly solvable model of disordered systems. *Phys. Rev. B*, 24 :2613–2626, 1981.
- [DGLR89] B. Diu, C. Guthmann, D. Lederer, and B. Roulet. *Physique Statistique*. Collection Enseignement des Sciences. Hermann, Paris, 1989.
- [DLL62] Martin Davis, George Logemann, and Donald Loveland. A machine program for theorem-proving. *Commun. ACM*, 5(7) :394–397, 1962.
- [DMU04] C. Di, A. Montanari, and R. Urbanke. Weight distributions of LDPC code ensembles : Combinatorics meets statistical physics. In *International Symposium on Information Theory*. IEEE, 2004.
- [DPTTJR02] C. Di, D. Proietti, I. E. Telatar, and R. L. Urbanke T. J. Richardson. Finite length analysis of low-density parity-check codes on the binary erasure channel. *IEEE Trans. Inform. Theory*, 48 :1570–1579, 2002.
- [DRU06] C. Di, T. J. Richardson, and R.L. Urbanke. Weight distribution of low-density parity-check codes. *IEEE Trans. Inform. Theory*, 52 :4839–4855, 2006.
- [Ell85] R. S. Ellis. *Entropy, Large Deviations, and Statistical Mechanics*. Springer-Verlag, New-York, 1985.
- [Ell95] R. S. Ellis. An overview of the theory of large deviations and applications to statistical physics. *Scand. Acturial J.*, 1 :97–142, 1995.
- [ER59] P. Erdős and A. Rényi. On random graphs. *Publ. Math. Debrecen*, 6 :290–297, 1959.
- [ER60] P. Erdős and A. Rényi. On the evolution of random graphs. *Publ. Math. Inst. Hungar. Acad. Sci.*, 5 :17–61, 1960.

- [ES02] D. J. Evans and D. J. Searles. The Fluctuation Theorem. *Advances in Physics*, 51 :1529–1585, November 2002.
- [FD07] B. J. Frey and D. Dueck. Clustering by Passing Messages Between Data Points. *Science*, 315 :972, 2007.
- [FL03] S. Franz and M. Leone. Replica bounds for optimization problems and diluted spin systems. *J. Stat. Phys.*, 3-4 :535–564, 2003.
- [FLMRT02] S. Franz, M. Leone, A. Montanari, and F. Ricci-Tersenghi. The dynamic phase transition for decoding algorithms. *Phys. Rev. E*, 66 :046120, 2002.
- [FP83] J. Franco and M. Paull. Probabilistic analysis of the davis-putnam procedure for solving satisfiability. *Discrete Applied Mathematics*, 5 :77–87, 1983.
- [Fri99] E. Friedgut. Sharp thresholds of graph properties, and the k -sat problem. *J. Amer. Math. Soc.*, 12, 1999.
- [Gal62] R. G. Gallager. Low-density parity check codes. *IRE Trans. Inf. Theory*, IT-8 :21, 1962.
- [Gal68] R. G. Gallager. *Information theory and reliable communication*. John Wiley and Sons, New York, 1968.
- [GJ79] M.R. Garey and D.S. Johnson. *Computers and intractability : a guide to the theory of NP-completeness*. Freeman, San Francisco, 1979.
- [GM84] D.J. Gross and M. Mézard. The simplest spin glass. *Nucl. Phys. B*, 240 :431, 1984.
- [Hay97] B. Hayes. Can’t get no satisfaction. *American scientist*, 85 :108–112, 1997.
- [JLR00] S. Janson, T. Luczak, and A. Rucinski. *Random graphs*. Wiley, New-York, 2000.
- [KFL01] F. R. Kschischang, B. Frey, and H.-A. Loeliger. Factor graphs and the sum-product algorithm. *IEEE Trans. Inform. Theory*, 47(2) :498–519, 2001.
- [KKKS98] Lefteris M. Kirousis, Evangelos Kranakis, Danny Krizanc, and Yanis C. Stamatiou. Approximating the unsatisfiability threshold of random formulas. *Random Structures and Algorithms*, 12(3) :253–269, 1998.
- [KM05] V. Kalapala and C. Moore. The Phase Transition in Exact Cover. *ArXiv Computer Science e-prints*, August 2005.
- [KMRT⁺07] Florent Krzakala, Andrea Montanari, Federico Ricci-Tersenghi, Guilhem Semerjian, and Lenka Zdeborova. Gibbs states and the set of solutions of random constraint satisfaction problems. *Proc. Natl. Acad. Sci.*, 104 :10318, 2007.

- [KR98] H. J. Kappen and F. B. Rodríguez. Efficient learning in boltzmann machines using linear response theory. *Neural Comput.*, 10(5) :1137–1156, 1998.
- [KS94] S. Kirkpatrick and B. Selman. Critical behavior in the satisfiability of random boolean expression. *Science*, 264 :1297–1301, 1994.
- [KS04] Y. Kabashima and D. Saad. Statistical mechanics of low-density parity-check codes. *J. Phys. A : Math. Gen.*, 37 :R1–R43, 2004.
- [KS05] K. M. Krishnan and P. Shankar. On the Complexity of finding Stopping Distance in Tanner Graphs. *ArXiv Computer Science e-prints*, December 2005.
- [KV03] N. Kashyap and A. Vardy. Stopping sets in codes from designs. In *Proc. Intern. Symp. on Inform. Theory (ISIT'03)*, page 122, 2003.
- [Lad75] Richard E. Ladner. On the structure of polynomial time reducibility. *J. ACM*, 22(1) :155–171, 1975.
- [LMS⁺97] Michael G. Luby, Michael Mitzenmacher, M. Amin Shokrollahi, Daniel A. Spielman, and Volker Stemann. Practical loss-resilient codes. In *STOC '97 : Proceedings of the twenty-ninth annual ACM symposium on Theory of computing*, pages 150–159, New York, NY, USA, 1997. ACM Press.
- [LMSS01] M. G. Luby, M. Mitzenmacher, M. Amin Shokrollahi, and D. A. Spielman. Efficient erasure correcting codes. *IEEE Transactions on Information Theory*, 47(2) :569–584, 2001.
- [LP86] L. Lovasz and M. D. Plummer. *Matching Theory*. North-Holland, Amsterdam, New York, 1986.
- [Mac99] D. J. C. MacKay. Good error-correcting codes based on very sparse matrices. *IEEE Trans. Inform. Theory*, 45(2) :399–431, 1999.
- [Mac03] D. J. C. MacKay. *Information theory, inference, and learning algorithms*. Cambridge University Press, Cambridge, 2003.
- [MM06a] Marc Mézard and Andrea Montanari. Reconstruction on trees and spin glass transition. *J. Stat. Phys.*, 124 :1317–1350, september 2006.
- [MM06b] T. Mora and M. Mézard. Geometrical organization of solutions to random linear Boolean equations. *Journal of Statistical Mechanics : Theory and Experiment*, 10 :P10007, October 2006.
- [MM07] M. Mézard and A. Montanari. *Constraint Satisfaction Networks in Physics and Computation*. 2007. En préparation, disponible sur www.lptms.u-psud.fr/membres/mezard/.
- [MMRU04] C. Measson, A. Montanari, T. Richardson, and R. Urbanke. Life above threshold : from list decoding to area theorem and MSE. In *Proc. ITW*, San Antonio, USA, October 2004.

- [MMS06] E. Marinari, R. Monasson, and G. Semerjian. An algorithm for counting circuits : Application to real-world and random graphs. *Europhysics Letters*, 73 :8–14, January 2006.
- [MMU05] C. Measson, A. Montanari, and R. Urbanke. Maxwell Construction : The Hidden Bridge between Iterative and Maximum a Posteriori Decoding. *ArXiv Computer Science e-prints*, June 2005.
- [MMW05] Elitza Maneva, Elchanan Mossel, and Martin J. Wainwright. A new look at survey propagation and its generalizations. In *SODA '05 : Proceedings of the sixteenth annual ACM-SIAM symposium on Discrete algorithms*, pages 1089–1098, Philadelphia, PA, USA, 2005. Society for Industrial and Applied Mathematics.
- [MMZ05a] M. Mézard, T. Mora, and R. Zecchina. Clustering of solutions in the random satisfiability problem. *Physical Review Letters*, 94 :197205, 2005.
- [MMZ05b] T. Mora, M. Mézard, and R. Zecchina. Pairs of sat assignments and clustering in random boolean formulae, 2005. cond-mat/0506053.
- [MMZ06] Stephan Mertens, Marc Mézard, and Riccardo Zecchina. Threshold values of random k-sat from the cavity method. *Random Struct. Algorithms*, 28(3) :340–373, 2006.
- [MN95] David J. C. MacKay and R. M. Neal. Good codes based on very sparse matrices. In *Proceedings of the 5th IMA Conference on Cryptography and Coding*, pages 100–111, London, UK, 1995. Springer-Verlag.
- [MN96] D. J. C. MacKay and R. M. Neal. Near Shannon limit performance of low density parity check codes. *Electronics Letters*, 32(18) :1645–1646, August 1996. Reprinted *Electronics Letters*, vol 33, no 6, 13th March 1997, p.457–458.
- [MO03] T. Murayama and M. Okada. One step RSB scheme for the rate distortion function. *Journal of Physics A Mathematical General*, 36 :11123–11130, October 2003.
- [Mon98] R. Monasson. Optimization problems and replica symmetry breaking in finite connectivity spin-glasses. *J. Phys. A*, 31 :515, 1998.
- [Mon01] A. Montanari. The glassy phase of Gallager codes. *Eur. Phys. J. B.*, 23 :121–136, 2001.
- [Mon05] A. Montanari. Tight bounds for ldpc and ldgm codes under map decoding. *IEEE Trans. Inform. Theory*, 51 :3221–3246, 2005.
- [MP86] M. Mézard and G. Parisi. Mean-field equations for the matching and the travelling salesman problem. *Europhys. Lett.*, 2 :913–918, 1986.

- [MP87] M. Mézard and G. Parisi. On the solution of the random link matching problems. *J. Physique*, 48 :1451–1459, 1987.
- [MP01] M. Mézard and G. Parisi. The bethe lattice spin glass revisited. *Eur. Phys. J. B*, 20 :217, 2001.
- [MP03] David J. C. MacKay and M. J. Postol. Weaknesses of Margulis and Ramanujan–Margulis low-density parity-check codes. In *Proceedings of MFCSIT2002, Galway*, volume 74 of *Electronic Notes in Theoretical Computer Science*. Elsevier, 2003.
- [MPR05] M. Mézard, M. Palassini, and O. Rivoire. Landscape of solutions in constraint satisfaction problems. *Phys. Rev. Lett.*, 95 :200202, 2005.
- [MPRT04] A. Montanari, G. Parisi, and F. Ricci-Tersenghi. Instability of one-step replica-symmetry-broken phase in satisfiability problems. *J. Phys. A*, 37 :2073, 2004.
- [MPV86] M. Mézard, G. Parisi, and M. A. Virasoro. SK model : the replica solution without replicas. *Europhysics Letters*, 1 :77, January 1986.
- [MPV87] M. Mézard, G. Parisi, and M. A. Virasoro. *Spin-Glass Theory and Beyond*, volume 9 of *Lecture Notes in Physics*. World Scientific, Singapore, 1987.
- [MPWZ02] R. Mulet, A. Pagnani, M. Weigt, and R. Zecchina. Coloring random graphs. *Phys. Rev. Lett.*, 89 :268701, 2002.
- [MPWZ07] Hamed Mahmoudi, Andrea Pagnani, Martin Weigt, and Riccardo Zecchina. Propagation of external regulation and asynchronous dynamics in random boolean networks, 2007.
- [MPZ02] M. Mézard, G. Parisi, and R. Zecchina. Analytic and algorithmic solution of random satisfiability problems. *Science*, 297 :812–815, 2002.
- [MR05] A. Montanari and T. Rizzo. How to compute loop corrections to the Bethe approximation. *Journal of Statistical Mechanics : Theory and Experiment*, 10 :P10011, October 2005.
- [MR06a] T. Mora and O. Rivoire. Error exponents of low-density parity-check codes on the binary erasure channel. In *Proc. ITW*, pages 81–85, Chengdu, China, october 2006.
- [MR06b] T. Mora and O. Rivoire. Statistical mechanics of error exponents for error-correcting codes. *Phys. Rev. E*, 74(5) :056110, November 2006.
- [MRT03] A. Montanari and F. Ricci-Tersenghi. On the nature of the low-temperature phase in discontinuous mean-field spin glasses. *Eur. Phys. J. B*, 33 :339, 2003.
- [MRTZ03] M. Mézard, F. Ricci-Tersenghi, and R. Zecchina. Alternative solutions to diluted p -spin models and XORSAT problems. *J. Stat. Phys.*, 111 :505, 2003.

- [MS05] A. Montanari and G. Semerjian. From large scale rearrangements to mode coupling phenomenology. *Phys. Rev. Lett.*, 94 :247201, 2005.
- [MS06a] E. Marinari and G. Semerjian. On the number of circuits in random graphs. *Journal of Statistical Mechanics : Theory and Experiment*, 6 :P06019, June 2006.
- [MS06b] A. Montanari and G. Semerjian. On the dynamics of the glass transition on bethe lattices. *J. Stat. Phys.*, 124 :103–189, 2006. cond-mat/0509366.
- [MT72] Raymond E. Miller and James W. Thatcher, editors. *Complexity of computer computations*. Plenum Press, New York, 1972.
- [MT06] Andrea Montanari and David Tse. Analysis of belief propagation for non-linear problems : The example of cdma (or : How to prove tanaka’s formula), 2006.
- [Mur04] T. Murayama. Thouless-Anderson-Palmer approach for lossy compression. *Phys. Rev. E*, 69(3) :035105, March 2004.
- [MW06] Emin Martinian and Martin Wainwright. Low density codes achieve therate-distortion bound. In *DCC ’06 : Proceedings of the Data Compression Conference (DCC’06)*, pages 153–162, Washington, DC, USA, 2006. IEEE Computer Society.
- [MZ96] R. Monasson and R. Zecchina. Entropy of the K-satisfiability problem. *Phys. Rev. Lett.*, 76 :3881–3885, 1996.
- [MZ97] Rémi Monasson and Riccardo Zecchina. Statistical mechanics of the random k -satisfiability model. *Phys. Rev. E*, 56(2) :1357–1370, Aug 1997.
- [MZ02] M. Mézard and R. Zecchina. Random k -satisfiability problem : From an analytic solution to an efficient algorithm. *Phys. Rev. E*, 66 :056126, 2002.
- [MZK⁺99] R. Monasson, R. Zecchina, S. Kirkpatrick, B. Selman, and L. Troyansky. Determining computational complexity from characteristic phase transitions. *Nature*, 400 :133–137, 1999.
- [New03] M. E. J. Newman. The structure and function of complex networks. *SIAM Review*, 45(2) :167–256, 2003.
- [Nis01] H. Nishimori. *Statistical Physics of Spin Glasses and Information Processing : An Introduction*. Oxford University Press, Oxford, UK, 2001.
- [OVZ05] A. Orlitsky, K. Viswanathan, and J. Zhang. Stopping set distribution of ldpc code ensembles. *IEEE Trans. Inform. Theory*, 51 :929–953, 2005.
- [PA87] C. Peterson and R. Anderson. A mean field theory learning algorithm for neural networks. *Complex Systems*, 1 :995–1019, 1987.

- [Pap91] Christos H. Papadimitriou. On selecting a satisfying truth assignment (extended abstract). In *Proceedings of the 32nd annual symposium on Foundations of computer science*, pages 163–169, Los Alamitos, CA, USA, 1991. IEEE Computer Society Press.
- [Pap94] C. H. Papadimitriou. *Computational complexity*. Addison-Wesley, 1994.
- [PSW96] B. Pittel, J. Spencer, and N.C. Wormald. Sudden emergence of a giant k -core in a random graph. *J. Comb. Theory Ser. B*, 67 :111–151, 1996.
- [PW06] Marco Pretti and Martin Weigt. Sudden emergence of q -regular subgraphs in random graphs. *Europhysics Letters*, 75 :8, 2006.
- [RB06] O. Rivoire and J. Barré. Exactly Solvable Models of Adaptive Networks. *Physical Review Letters*, 97(14) :148701, October 2006.
- [RBMM04] O. Rivoire, G. Biroli, O. C. Martin, and M. Mézard. Glass models on bethe lattices. *Eur. Phys. J. B*, 37 :55–78, 2004.
- [Riv04] O. Rivoire. Properties of atypical graphs from negative complexities. *J. Stat. Phys.*, 117 :453, 2004.
- [Riv05] O. Rivoire. The cavity method for large deviations. *J. Stat. Mech.*, page P07004, 2005.
- [RSU01] T. J. Richardson, M. A. Shokrollahi, and R. L. Urbanke. Design of capacity-approaching irregular low-density parity-check codes. *IEEE Trans. Inform. Theory*, 47 :619–637, 2001.
- [RSZ07] J. Raymond, A. Sportiello, and L. Zdeborová. The Phase Diagram of 1-in-3 Satisfiability Problem. *ArXiv Condensed Matter e-prints*, February 2007.
- [RU01] Richardson and Urbanke. The capacity of low-density parity-check codes under message-passing decoding. *IEEE Trans. Inform. Theory*, 47, 2001.
- [RU07] T. Richardson and R. Urbanke. *Modern Coding Theory*. 2007. En préparation, disponible sur lthcwww.epfl.ch/mct.
- [SBSB06] E. Schneidman, M. J. Berry, R. Segev, and W. Bialek. Weak pairwise correlations imply strongly correlated network states in a neural population. *Nature*, 440 :1007–1012, April 2006.
- [SCCV05] M. G. Stepanov, V. Chernyak, M. Chertkov, and B. Vasic. Diagnosis of Weaknesses in Modern Error Correction Codes : A Physics Approach. *Physical Review Letters*, 95(22) :228701, November 2005.
- [Sch78] T. J. Schaefer. The complexity of satisfiability problems. In *Proc. 10th STOC*, page 216, San Diego, CA, USA, 1978. ACM.

- [Sem07] Guilhem Semerjian. On the freezing of variables in random constraint satisfaction problems, 2007. Preprint arXiv.org :0705.2147.
- [Sha48] C. E. Shannon. A mathematical theory of communication. *Bell System Tech. Journal*, 27 :379–423, 623–655, 1948.
- [SK75] D. Sherrington and S. Kirkpatrick. Solvable model of a spin-glass. *Phys. Rev. Lett.*, 35 :1792–1796, 1975.
- [SM03] Guilhem Semerjian and Rémi Monasson. Relaxation and metastability in a local search procedure for the random satisfiability problem. *Phys. Rev. E*, 67(6) :066103, Jun 2003.
- [SM04] G. Semerjian and R. Monasson. A study of pure random walk on random satisfiability problems with “physical” methods. In E. Giunchiglia and A. Tachella, editors, *Proceedings of the SAT 2003 conference*, volume 120 of *Lecture Notes in Computer Science*, page 2919. Springer, 2004.
- [SML96] Bart Selman, David G. Mitchell, and Hector J. Levesque. Generating hard satisfiability problems. *Artif. Intell.*, 81(1-2) :17–29, 1996.
- [Sou89] N. Surlas. Spin-glass models as error-correcting codes. *Nature*, 339 :693–694, 1989.
- [Sou94] N. Surlas. Spin-glasses, error-correcting codes and finite-temperature decoding. *Europhys. Lett.*, 25 :159–164, 1994.
- [Tal00] M. Talagrand. Rigorous low temperature results for the p-spin mean field spin glass model. *Probability Theory and Related Fields*, 117 :303–360, 2000.
- [Tal03] M. Talagrand. *Spin glasses : a challenge for mathematicians. Cavity and mean field models*. Springer-Verlag, New-York, 2003.
- [Tal06] M. Talagrand. The parisi formula. *Ann. Math.*, 163 :221–263, 2006.
- [Tan81] Robert Michael Tanner. A recursive approach to low complexity codes. *IEEE Transactions on Information Theory*, 27(5) :533–547, 1981.
- [Tan98] T. Tanaka. Mean-field theory of Boltzmann machine learning. *Phys. Rev. E*, 58 :2302–2310, August 1998.
- [TSBB06] G. Tkacik, E. Schneidman, M. J. I. Berry, and W. Bialek. Ising models for networks of real neurons. *eprint arXiv :q-bio/0611072*, November 2006.
- [Var97] A. Vardy. The intractability of computing the minimum distance of a code. *IEEE Trans. Inf. Theory*, 43 :1757–1766, 1997.
- [WH00] M. Weigt and A. K. Hartmann. The number of guards needed by a museum : A phase transition in vertex covering of random graphs. *Phys. Rev. Lett.*, 84 :6118, 2000.

- [WH01] M. Weigt and A. K. Hartmann. Minimal vertex covers on finite-connectivity random graphs : A hard-sphere lattice-gas picture. *Phys. Rev. E*, 63 :056127, 2001.
- [WM03] M.J. Wainwright and E. Maneva. Lossy source encoding via message-passing and decimation over generalized codewords of ldgm codes. In *Proc. Intern. Symp. on Inform. Theory (ISIT'03)*, pages 1493–1497, 2003.
- [Yed01] Jonathan Yedidia. An idiosyncratic journey beyond mean field theory. In Manfred Opper and David Saad, editors, *Advanced Mean Field Methods, Theory and Practice*, pages 21–36. The MIT Press, 2001.
- [YFW02] J. S. Yedidia, W. F. Freeman, and Y. Weiss. Constructing free energy approximations and generalized belief propagation algorithms. *technical report TR-2002-35, Mitsubishi Electrical Research Laboratories*, 2002. available at <http://www.merl.com>.
- [Zho05] H. Zhou. Long-Range Frustration in a Spin-Glass Model of the Vertex-Cover Problem. *Physical Review Letters*, 94(21) :217203, June 2005.
- [ZK] Lenka Zdeborova and Florent Krzakala. Phase transitions in the coloring of random graphs. Preprint arXiv.org :0704.1269.
- [ZM06] L. Zdeborová and M. Mézard. The number of matchings in random graphs. *Journal of Statistical Mechanics : Theory and Experiment*, 5 :P05003, May 2006.

Résumé

Les problèmes d'optimisation et de satisfaction de contraintes sur des ensembles de variables discrètes sont l'objet principal de la complexité algorithmique. Ces problèmes ont récemment bénéficié des outils et des concepts de la physique des systèmes désordonnés, à la fois théoriquement et algorithmiquement. En particulier, il a été suggéré que les difficultés pratiques soulevées par certaines instances dures de problèmes d'optimisation sont liées à la structure fragmentée de leur espace de solutions, qui rappelle une phase vitreuse. Parallèlement, les codes de correction d'erreur de pointe, qui peuvent être ramenés à des problèmes d'optimisation, reposent sur la séparabilité de leurs messages afin d'assurer une communication fiable. L'objet de cette thèse est d'explorer, dans un cadre commun, cette relation entre les propriétés d'inférence et l'organisation géométrique, dans les problèmes issus de la complexité algorithmique et de la théorie de l'information.

Après une introduction physique des problèmes et des concepts liés aux domaines sus-évoqués, les méthodes de passage de messages, basées sur l'approximation de Bethe, sont introduites. Ces méthodes sont utiles d'un point de vue physique, car elle permettent d'étudier les propriétés thermodynamiques d'ensemble d'instances aléatoires. Elles sont également utiles pour l'inférence. L'analyse de spectres de distances est ensuite effectuée à l'aide de méthodes combinatoires et de passage de messages, et mises à profit afin de prouver et l'existence de la fragmentation dans les problèmes de satisfaction de contraintes, et d'en étudier les aspects importants.