

Transmission robuste de zéro-forçage asymptotiquement optimale pour coopération imparfaite de transmetteurs

Antonio BAZCO-NOGUERAS^{1,2}, Lorenzo MIRETTI¹, Paul DE KERRET¹, David GESBERT¹, Nicolas GRESSET²

¹Département Systèmes de Communication, EURECOM
Campus SophiaTech, 450 Route des Chappes, 06410 Biot, France

²Mitsubishi Electric R&D Centre Europe (MERCE)
1, allée de Beaulieu, 35700 Rennes, France

{bazco,miretti,dekerret,gesbert}@eurecom.fr, n.gresset@fr.mercedes-mee.com

Résumé – Dans cet article*, nous étudions la transmission sans-fil de 2 transmetteurs servant conjointement 2 utilisateurs et disposant chacun d’une information de canal imparfaite et potentiellement différente de celle de l’autre transmetteur. Dans cette configuration, il est connu que les schémas de précodage conçus pour un système centralisé –où tous les transmetteurs partagent la même information de canal– souffrent d’une importante perte de performance suite aux différences entre les estimations des transmetteurs. Dans cette contribution, nous étendons les schémas de transmission robuste proposés dans les travaux précédents au cas où chaque canal peut être connu avec une qualité différente à chaque transmetteur, et nous mettons en évidence des configurations d’information distribuée pour lesquelles il est possible d’atteindre asymptotiquement le même débit total que dans le cas centralisé.

Abstract – In this work* we analyze a Network MIMO channel with 2 Transmitters (TXs) jointly serving 2 users (RXs), where each TX has a different multi-user Channel State Information (CSI), potentially with a different accuracy at each one of the links TX-RX. It was recently shown that this distributed scenario does reach out to the same Degrees-of-Freedom –or multiplexing gain– as its genie-aided centralized counterpart in which both TXs share the best-quality estimate of each link TX-RX. Nevertheless, the DoF analysis does not provide the actual rate and the rate gap between the centralized and the decentralized settings remained unknown. In this paper, we show that it is possible to achieve asymptotically the same sum rate as that attained by Zero-Forcing precoding in a centralized setting endowed with the best-quality CSI, for any possible CSI allocation at the TXs. This is accomplished thanks to a new precoding scheme that adapts to the CSI allocation.

1 Introduction

La coopération entre transmetteurs (TXs) dans un canal de diffusion sans-fil est un sujet qui a été largement étudié dans la littérature. De ce fait, il est connu que, si les TXs ont une information du canal parfaite, la capacité du système est proportionnelle au nombre d’antennes [1]. Néanmoins, ce gain est fortement conditionné par la qualité de l’information de canal (CSI) aux TXs [1, 2], d’où l’importance de quantifier cette dépendance. Même si de nombreux travaux se sont portés sur cette problématique, ils font l’hypothèse que l’information de canal, bien qu’imparfaite, est *parfaitement* partagée entre tous les TXs. Cependant, cette hypothèse n’est pas valide pour un certain nombre de configurations importantes, comme par exemple lorsque les TXs sont séparés géographiquement et une partie d’eux a une liaison sans-fil limitée avec le réseau cœur.

Nous étudions ici le cas où deux TXs ayant chacun reçu une estimation de précision différente servent conjointement deux utilisateurs. La spécificité des résultats présentés par la suite provient du fait que chaque TX peut avoir une précision différente *pour chaque coefficient de canal*. Il a été récemment mis en évidence

qu’il était possible d’atteindre le même facteur pré-logarithmique –aussi appelé « Degrés-de-Liberté » (DoF)– que le scénario où la meilleure estimation pour chaque coefficient du canal est mise à disposition de tous les TXs [3]. Ce résultat est atteint à travers une méthode de précodage qui adapte l’usage que chaque TX fait de son estimée à la configuration globale d’information de canal [3].

Cependant, le DoF n’offre qu’une information incomplète sur le débit asymptotique, et il est important de connaître ce débit avec plus de précision. Dans le cas particulier où un TX a une meilleure précision pour tous les coefficients de canal, il a été montré qu’il est possible d’atteindre à haut rapport signal-à-bruit le même débit que dans le cas centralisé se basant sur l’information de meilleure qualité [4]. Ce résultat surprenant montre qu’il est possible de réduire fortement les pertes dues au partage imparfait de l’information entre TXs en adaptant la méthode de transmission à cette dégradation.

La contribution principale de ce travail consiste en l’extension de l’approche de [4] à toutes les autres configurations d’information de canal. Ce résultat est obtenu en combinant l’approche de [4] avec la méthode de transmission décrite dans [3]. Ainsi, nous mettons en évidence que le débit du scénario centralisé utilisant le zéro-forçage est asymptotiquement atteint, indépendamment de quel TX dispose de l’estimation la plus précise pour un coefficient donné.

*L. Miretti, P. de Kerret, et D. Gesbert sont soutenus par l’European Research Council sous le programme de recherche et d’innovation de l’Union européenne Horizon 2020 (Accord no. 670896).

2 Description du scénario

2.1 Modèle de transmission

Nous étudions une transmission dans un canal sans-fil où 2 TXs avec d'information imparfaite servent conjointement 2 récepteurs (RXs) et où les quatre nœuds n'ont qu'une seule antenne chacun. Le signal reçu par le i -ème RX peut alors s'écrire :

$$y_i = \mathbf{h}_i^H \mathbf{x} + z_i \quad (1)$$

où $\mathbf{h}_i^H \in \mathbb{C}^{1 \times 2}$ dénote le canal vers le RX i , $\mathbf{x} \in \mathbb{C}^{2 \times 1}$ est le signal multi-utilisateur transmis, et $z_i \in \mathbb{C}$ est le bruit additif au RX i , distribué comme $\mathcal{N}_{\mathbb{C}}(0,1)$. Nous dénotons également le canal entre le TX k et le RX i comme $h_{i,k}$, et le canal multi-utilisateur comme $\mathbf{H} \triangleq [\mathbf{h}_1, \mathbf{h}_2]^H \in \mathbb{C}^{2 \times 2}$, de telle façon que $h_{i,k} = \{\mathbf{H}\}_{i,k}$. Le canal est obtenu à partir d'une distribution générique de telle sorte que toutes les matrices de canal et leurs sous-matrices sont presque sûrement de rang plein. Du fait qu'on considère un scénario 2×2 , tous les indices utilisés dans ce papier appartiennent à l'ensemble $\{1, 2\}$. On omettra donc de mentionner cet ensemble lorsque ceci ne crée pas d'ambiguïté. Le signal multi-utilisateur émis $\mathbf{x} \in \mathbb{C}^{2 \times 1}$ est obtenu à partir du précodage des données utilisateurs $\mathbf{s} \triangleq [s_1 \ s_2]^T$, où les symboles s_i sont i.i.d. selon la loi $\mathcal{N}_{\mathbb{C}}(0,1)$ et s_i désigne le symbole à transmettre au RX i , de telle manière que

$$\mathbf{x} \triangleq \bar{P} \begin{bmatrix} \mathbf{t}_1 & \mathbf{t}_2 \end{bmatrix} \begin{bmatrix} s_1 \\ s_2 \end{bmatrix}, \quad (2)$$

où $\bar{P} \triangleq \sqrt{P}$ et P est la puissance maximale transmise par TX. Le vecteur $\mathbf{t}_i \in \mathbb{C}^{2 \times 1}$ désigne le vecteur normalisé de précodage pour les données destinées au RX i . Nous définissons aussi le précodeur multi-utilisateur $\mathbf{T}$ et le vecteur de précodage appliqué au TX j $\mathbf{t}_{\text{TX}j}$ de telle façon que

$$\mathbf{T} = \begin{bmatrix} \mathbf{t}_1 & \mathbf{t}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{t}_{\text{TX}1} \\ \mathbf{t}_{\text{TX}2} \end{bmatrix}. \quad (3)$$

De manière importante, nous considérons à chaque TX une contrainte de puissance *instantanée* pour le précodeur, i.e., $\forall j$,

$$\|\mathbf{t}_{\text{TX}j}\|_2 \leq 1. \quad (4)$$

Dans le cas distribué, cette contrainte de puissance devient très importante car la limite de puissance à un TX impacte la décision à l'autre TX.

2.2 Information de canal distribuée

Les estimées imparfaites de canal disponibles aux TXs sont définies de manière similaire au canal : L'estimation du canal multi-utilisateur obtenue par le TX j est écrite $\hat{\mathbf{H}}^{(j)}$, alors que $\hat{h}_{i,k}^{(j)}$ dénote l'estimation pour le canal entre TX k et RX i . Dans le cas d'une information de canal distribuée, la précision à chaque $\hat{h}_{i,k}^{(j)}$ est différente pour chaque i, j, k . Plus précisément, on modèle l'erreur d'estimation comme

$$\hat{h}_{i,k}^{(j)} \triangleq z_{i,k}^{(j)} h_{i,k} + z_{i,k}^{(j)} \delta_{i,k}^{(j)}, \quad (5)$$

où $z_{i,k}^{(j)} \triangleq \sqrt{1 - (z_{i,k}^{(j)})^2}$, $z_{i,k}^{(j)}$ est une variable aléatoire réelle qui détermine la variance du bruit d'estimation, et $\delta_{i,k}^{(j)} \in \mathbb{C}$ est une variable aléatoire de moyenne nulle et de variance unitaire. La variance du bruit est modélisée de telle façon qu'elle décroît lorsque le rapport signal-à-bruit P augmente :

$$\mathbb{E}[(z_{i,k}^{(j)})^2] = O(P^{-\alpha_{i,k}^{(j)}}), \quad (6)$$

où $\alpha_{i,k}^{(j)}$ est appelé le *coefficient de qualité de l'information de canal* au TX j . Intuitivement, $\alpha = 0$ correspond à une information de canal inexploitable, tandis que $\alpha = 1$ correspond à une information de canal « parfaite » dans le sens où les erreurs d'estimation deviennent asymptotiquement négligeables [1, 2]. Dans ce document, nous considérons que $\alpha_{i,k}^{(j)} > 0 \ \forall i, j, k$. Les $\alpha_{i,k}^{(j)}$ dépendent de la topologie du réseau et varient donc lentement. En conséquence, on suppose que tous les coefficients peuvent être connus parfaitement par tous les TXs. De plus, il faut bien noter que TX j calcule ses coefficients de précodage à partir de $\hat{\mathbf{H}}^{(j)}$ *sans avoir la possibilité d'échanger d'autres informations avec l'autre TX*.

2.3 Précodage de zéro-forçage centralisé

Comme étape préliminaire, nous commençons par présenter le précodage centralisé qui servira de référence pour l'évaluation des performances. Nous nous concentrons sur les méthodes de transmission utilisant le « Zéro-Forçage » (ZF) des interférences –c'est-à-dire où l'objectif du précodeur est d'annuler les interférences reçues aux RXs– car elles seules permettent une compréhension analytique. Autrement dit, si $\mathbf{v}_i$ est un vecteur de précodage ZF de puissance unitaire pour les données de RX i , et $\hat{\mathbf{h}}_i^H$ est l'estimation du vecteur de canal en direction de l'autre RX, le vecteur $\mathbf{v}_i$ doit vérifier $\hat{\mathbf{h}}_i^H \mathbf{v}_i = 0$. On peut alors écrire $\mathbf{T}^{\text{ZF}} \triangleq [\mu_1 \mathbf{v}_1 \ \mu_2 \mathbf{v}_2]$, μ_i étant un paramètre dessiné pour satisfaire la limite de puissance dans (4).

Afin d'analyser la performance du schéma proposé, et d'étudier quelle est la perte de débit due à la nature distribuée de l'information du canal, nous comparerons le débit moyen obtenu par notre schéma dans le cas distribué où la précision d'estimation est donnée par l'ensemble $\{\alpha_{i,k}^{(1)}, \alpha_{i,k}^{(2)}\}_{\forall i,k}$ (dénnoté par R_{distr}) avec le débit moyen obtenu dans le cas « idéal » où les deux TXs partagent la meilleure estimation pour chaque lien (dénnoté par R_{centr}), i.e., ayant la précision

$$\{\alpha_{i,k}^{\text{id}} = \max(\alpha_{i,k}^{(1)}, \alpha_{i,k}^{(2)})\}_{\forall i,k}. \quad (7)$$

2.4 Débit moyen total

La métrique utilisée est le « débit moyen total ». Le débit moyen pour RX i est donné par

$$R_i \triangleq \mathbb{E} \left[\log_2 \left(1 + \frac{P |\mathbf{h}_i^H \mathbf{t}_i|^2}{1 + P |\mathbf{h}_i^H \mathbf{t}_j|^2} \right) \right] \quad (8)$$

et le « débit moyen total » est donc écrit $R \triangleq R_1 + R_2$.

3 Schéma proposé

Avant de décrire notre principal résultat dans la Section 4, nous commençons par présenter le schéma de transmission utilisé.

a) ZF décentralisé :

De manière similaire au schéma centralisé, dans le cas distribué on peut diviser le vecteur de précodage pour RX i ($\mathbf{t}_i$) en un vecteur qui prend en charge le zéro-forçage d'interférence ($\mathbf{w}_i$) et un paramètre permettant de satisfaire la limite de puissance à chaque TX (λ_i). Cependant, dans le cas distribué, chaque TX j choisit son coefficient $\{\mathbf{t}_i\}_j$ uniquement à partir de son estimée locale, de telle sorte que le précodeur effectif appliqué aux données de RX i s'écrit :

$$\mathbf{t}_i = [\lambda_i^{(1)} \{\mathbf{w}_i^{(1)}\}_1 \quad \lambda_i^{(2)} \{\mathbf{w}_i^{(2)}\}_2]^T \quad (9)$$

TABLE 1 – Précodeur $\mathbf{w}_i = [\{\mathbf{w}_i^{(1)}\}_1 \{\mathbf{w}_i^{(2)}\}_2]^T$ pour les données destinées au RX i en fonction de la répartition de CSI.

	CSI à TX 1	CSI Local	CSI Opposé
	$\alpha_{i,1}^{(1)} > \alpha_{i,1}^{(2)}$ $\alpha_{i,2}^{(1)} \geq \alpha_{i,2}^{(2)}$	$\alpha_{i,1}^{(1)} > \alpha_{i,1}^{(2)}$ $\alpha_{i,2}^{(1)} < \alpha_{i,2}^{(2)}$	$\alpha_{i,1}^{(1)} \leq \alpha_{i,1}^{(2)}$ $\alpha_{i,2}^{(1)} \geq \alpha_{i,2}^{(2)}$
$\mathbf{w}_i$	$\begin{bmatrix} (\hat{\mathbf{h}}_{i,1}^{(1)})^{-1} \hat{\mathbf{h}}_{i,2}^{(1)} \\ -1 \end{bmatrix}$	$\begin{bmatrix} (\hat{\mathbf{h}}_{i,1}^{(1)})^{-1} \\ -(\hat{\mathbf{h}}_{i,2}^{(2)})^{-1} \end{bmatrix}$	$\begin{bmatrix} \hat{\mathbf{h}}_{i,2}^{(1)} \\ -\hat{\mathbf{h}}_{i,1}^{(2)} \end{bmatrix}$

où l'indice en exposant (j) représente le fait que le coefficient est estimé au TX j .

Sachant que l'information est distribuée de manière hétérogène entre les TXs, on peut effectuer le précodage de telle façon que TX j utilise l'estimation $\hat{\mathbf{h}}_{i,k}^{(j)}$ pour calculer son coefficient de suppression d'interférence $\{\mathbf{w}_i\}_j$ *seulement* s'il est le TX le mieux informé sur le lien $\mathbf{h}_{i,k}$. De ce fait, on peut caractériser trois types de précodage en fonction de quel TX a la meilleure estimation de chaque lien. Les différents schémas de précodage sont présentés dans la Table 1. Nous avons omis le cas où TX 2 a la meilleure qualité d'estimation pour la matrice de canal entière car il peut être obtenu par symétrie en inversant le rôle des TXs dans la première colonne de la Table 1.

Il est important d'observer que le vecteur $\mathbf{w}_i$ dépend seulement des coefficients de meilleure qualité dans les trois régimes, ce qui est la clé pour atteindre le même niveau de réduction d'interférence que dans le cas centralisé utilisant l'information la plus précise. A l'opposé, λ_i dépend de la norme du vecteur $\mathbf{w}_i$ et est donc dépendant de toutes les estimations, les plus précises et les moins précises.

b) Quantification du paramètre de puissance :

Bien que le précodeur présenté dans Table 1 permette d'améliorer la suppression d'interférence par rapport au précodage centralisé naïf –i.e., celui qui n'est pas conscient de la nature distribuée de l'information–, la réduction d'interférence est encore limitée par les différences entre $\lambda_i^{(1)}$ et $\lambda_i^{(2)}$. Nous proposons dans ce travail de résoudre ce problème à travers la quantification de λ_i à un niveau de quantification adéquat. Si les deux TXs obtiennent la même valeur quantifiée, le précodeur s'écrit :

$$\mathbf{T}_i^{\text{Distr}} = [\lambda_1^Q \mathbf{w}_1 \quad \lambda_2^Q \mathbf{w}_2] \quad (10)$$

où $\lambda_i^Q \triangleq Q(\lambda_i^{(j)})$, $\forall j$ et $Q(\cdot)$ désigne le quantificateur. Dans ce cas, le précodage proposé permet de réaliser le même niveau d'atténuation de l'interférence que le schéma centralisé basé sur l'estimée de meilleure qualité. En revanche, l'utilisation d'une quantification mène à une granularité plus faible du contrôle de puissance et ainsi à une perte de puissance du signal désiré. Il faut donc choisir le niveau de quantification (et le quantificateur) de manière à réaliser un compromis entre la probabilité d'obtenir la même valeur aux deux TXs (et ainsi une bonne atténuation d'interférence) et un contrôle de puissance précis (et ainsi un signal désiré puissant). L'optimisation du quantificateur sera étudiée dans les prochains travaux et on considère par la suite un quantificateur uniforme, dont le pas de quantification peut être optimisé.

4 Résultat principal

Nous pouvons maintenant formuler le résultat théorique principal.

Proposition 1 *Pour tout $\alpha_{i,k}^{(j)} > 0$, le schéma décrit dans la Section 3 atteint un débit R_{distr} tel que*

$$\lim_{P \rightarrow \infty} R_{\text{centr}} - R_{\text{distr}} \leq 0, \quad (11)$$

où R_{centr} est le débit obtenu dans la configuration centralisée où les deux TXs ont accès à l'estimation la plus précise, i.e., où

$$\alpha_{i,k}^{\text{centr}} = \max(\alpha_{i,k}^{(1)}, \alpha_{i,k}^{(2)}). \quad (12)$$

Esquisse de preuve: Par soucis de brièveté, nous présentons seulement une esquisse de la démonstration afin de mettre en évidence les points les plus importants de la preuve. Celle-ci est en effet longue, mais très similaire à la preuve utilisée dans le travail [4]. Nous commençons par définir l'événement $\mathcal{A}$ correspondant aux réalisations pour lesquelles les valeurs quantifiées aux deux TXs sont égales (tout en excluant les cas dégénérés), c'est-à-dire

$$\mathcal{A} \triangleq \{(\hat{\mathbf{H}}^{(1)}, \hat{\mathbf{H}}^{(2)}) \mid \forall i, Q(\lambda_i^{(1)}) = Q(\lambda_i^{(2)}) \in \mathbb{R}^+\}. \quad (13)$$

Les principales étapes de la démonstration sont donc :

(i) Procédant comme [4, Section IV], on peut écrire :

$$R_{\text{centr}} - R_{\text{distr}} \leq \Gamma_1^{\text{AV}} + \Gamma_2^{\text{AV}} + \Pr(\mathcal{A}^c) \log_2(1+2P) \quad (14)$$

où $\Gamma_i^{\text{AV}} \triangleq \mathbb{E}_{|\Omega|} [\log_2(|\lambda_i^* / \lambda_i^Q|^2)]$ avec λ_i^* étant le paramètre obtenu par le schéma centralisé basé sur la meilleure estimation de chaque coefficient de canal. En d'autres termes, Γ_i^{AV} est une mesure de la différence entre la valeur centralisée λ_i^* et son homologue distribué λ_i^Q , lorsque les TXs utilisent la même valeur.

(ii) Le choix du quantificateur est fondamental pour atteindre (11). On met maintenant en évidence les propriétés importantes du quantificateur.

Lemma 1 *Soit $\mathcal{Q}_u$ un quantificateur uniforme avec un pas de quantification $q = \bar{P}^{-\frac{\alpha_m}{2}}$, où $\alpha_m = \min_{i,j,k}(\alpha_{i,k}^{(j)})$, tel que*

$$\mathcal{Q}_u(x) \triangleq \bar{P}^{-\frac{\alpha_m}{2}} \lfloor \bar{P}^{\frac{\alpha_m}{2}} x \rfloor. \quad (15)$$

Alors, $\mathcal{Q}_u$ satisfait la propriété

$$\lim_{P \rightarrow \infty} Q(\lambda_i^{(j)}) = \lambda_i^* \quad \text{p.s.} \quad \forall i, j, \quad (P1)$$

où p.s. signifie presque sûrement, ainsi que la propriété

$$\Pr(\mathcal{A}^c) = o\left(\frac{1}{\log_2(P)}\right). \quad (P2)$$

Esquisse de preuve: (Lemma 1) La propriété (P1) est vérifiée du fait que $\hat{\mathbf{H}}^{(j)}$ converge vers $\mathbf{H}$ p.s. lorsque $\alpha_{i,k}^{(j)} > 0$ et du fait que le pas de quantification tende vers 0. Pour prouver (P2), il faut exploiter le fait que $q = \bar{P}^{-\frac{\alpha_m}{2}}$ croît exponentiellement plus lentement que la pire des puissances $\bar{P}^{-\alpha_m}$. Ainsi la probabilité que l'erreur d'estimation mène à un désaccord dans la valeur quantifiée aux deux TXs peut-être bornée de manière à obtenir la propriété (P2). $\square$

(iii) La condition (P1) assure que $\lim_{P \rightarrow \infty} \Gamma_i^{\text{AV}} = 0$ alors que la condition (P2) assure que $\lim_{P \rightarrow \infty} \Pr(\mathcal{A}^c) \log_2(1+2P) = 0$. Pris ensemble, ces résultats permettent de conclure la preuve. $\square$

Le schéma proposé est particulièrement intéressant car il permet d'atteindre en même temps deux objectifs opposés : le DoF maximum et un contrôle de puissance asymptotiquement optimal. L'extension de ce résultat aux cas avec plusieurs antennes ou plus de RXs est actuellement étudiée par le groupe. Il est conjecturé que (11) soit atteint pour tous les cas où le TX avec l'information la plus précise aie un nombre d'antennes au moins égal au nombre de RXs interférés.

4.1 Simulations

Afin d'illustrer les résultats précédents, nous présentons des résultats de simulations numériques dans un scénario avec deux niveaux de précision différents : une estimation « précise » et une autre « approximative » pour chaque coefficient du canal. De cette façon, on sélectionne $\alpha_{i,k}^{\max} = 1$ pour l'estimation précise et $\alpha_{i,k}^{\min} = 0.4$ pour l'estimation approximative. Les résultats sont obtenus par moyennage sur 5000 réalisations de Monte-Carlo du canal sans-fil. Concernant le quantificateur Q , le pas de quantification est optimisé en utilisant une recherche exhaustive du pas optimale et en évaluant les performances avec des simulations de Monte-Carlo.

Comme illustre Table 1, il y a six cas différents pour les deux RXs. Concentrons-nous sur le cas où l'information de canal du RX 1 est distribuée de manière locale et celle de RX 2 de manière opposée :

$$\begin{array}{ccccc} \alpha_{1,1}^{(1)} = 1 & \alpha_{1,2}^{(1)} = 0.4 & \parallel & \alpha_{1,1}^{(2)} = 0.4 & \alpha_{1,2}^{(2)} = 1 \\ & & \parallel & & \\ \alpha_{2,1}^{(1)} = 0.4 & \alpha_{2,2}^{(1)} = 1 & \parallel & \alpha_{2,1}^{(2)} = 0.4 & \alpha_{2,2}^{(2)} = 1 \end{array}$$

Dans la Fig. 1, le débit atteint par le schéma proposé en fonction du SNR est comparé avec le débit atteint dans le cas centralisé utilisant l'estimée la plus précise (« ZF Centralisé idéal »), avec le zéro-forçage classique appliqué à chaque TX sans prendre en compte la configuration d'information de canal (« Naïf ZF »), et enfin avec le TDMA. On constate que le schéma proposé offre une performance supérieure aux autres schémas. De plus, il converge vers le cas idéal où les deux TXs partagent la meilleure estimation.

Dans le but d'illustrer le résultat de convergence de la Proposition 1, on présente dans la Fig. 2 le pourcentage du débit du ZF centralisé idéal qui est atteint par le schéma de transmission proposé pour toutes les configurations d'information de canal possibles. On montre donc côte à côte le pourcentage atteint à $P = 30, 50$ et 80 dB. Il est important de garder à l'esprit que les valeurs sont normalisées, bien que le débit centralisé de référence soit de 15 bits/Hz/s à 30 dB, 28 bits/Hz/s à 50 dB, et 48 bits/Hz/s à 80 dB, ce que est montré dans Fig. 1. De plus, il est à noter que le schéma naïf atteint un pourcentage de plus en plus petit à mesure que le SNR augmente.

5 Conclusions

Nous avons présenté un schéma de précodage linéaire permettant d'améliorer considérablement le débit lorsque les estimées de canal disponibles aux TXs sont imparfaites et imparfaitement partagées. La méthode de transmission proposée s'adapte à toutes les configurations d'information de canal et permet d'atteindre asymptotiquement les mêmes performances que le cas centralisé disposant de l'information la plus précise.

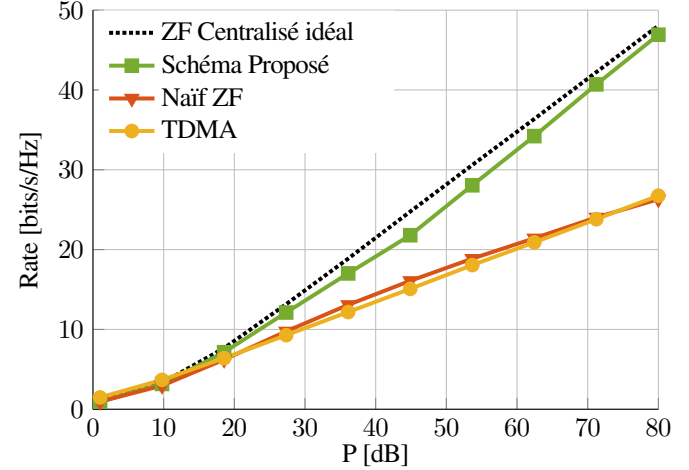


FIGURE 1 – Débit total pour le cas où l'information de canal du RX 1 est « locale » et celle de RX 2 est « opposée ».

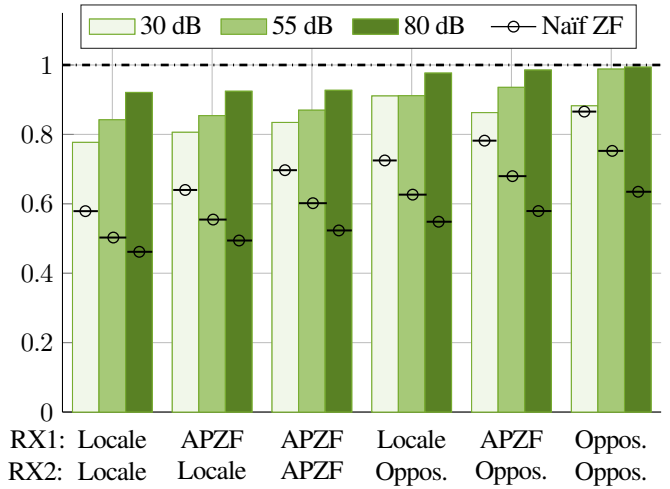


FIGURE 2 – Pourcentage du débit centralisé « idéal » atteint par le schéma proposé et par zéro-forçage naïf pour $P \in \{30, 55, 80\}$ dB.

Références

- [1] A. G. Davoodi and S. A. Jafar, “Aligned Image Sets Under Channel Uncertainty : Settling Conjectures on the Collapse of Degrees of Freedom Under Finite Precision CSIT,” *IEEE Trans. Inf. Theory*, vol. 62, no. 10, pp. 5603–5618, Oct 2016.
- [2] N. Jindal, “MIMO Broadcast Channels with finite-rate feedback,” *IEEE Trans. Inf. Theory*, vol. 52, no. 11, pp. 5045–5060, Nov. 2006.
- [3] A. Bazco-Nogueras, P. de Kerret, D. Gesbert, and N. Gresset, “Distributed CSIT Does Not Reduce the Generalized DoF of the 2-user MISO Broadcast Channel,” *IEEE Wireless Commun. Letters*, 2018, (Early access).
- [4] A. Bazco-Nogueras, L. Miretti, P. de Kerret, D. Gesbert, and N. Gresset, “Achieving Vanishing Rate Loss in Decentralized Network MIMO,” 2019. [Online]. Available : <https://arxiv.org/abs/1901.06849>