

Tehnica de Acces Multiplu prin Diviziune în Cod (CDMA)

- modulațiile studiate anterior căutau să asigure o eficiență spectrală cât mai ridicată în condițiile utilizării unei lărgimi de bandă cât mai reduse, pentru un debit binar impus.
- în cazul acestor tipuri de transmisii accesul utilizatorilor, în sistemele care deservesc mai mulți utilizatori, la banda de frecvență a sistemului se face fie prin diviziune în timp (TDMA), fie prin diviziune în frecvență (FDMA), fie prin combinații între cele două.
- de aceea în aceste sisteme se gestionează resurse de tip timp-frecvență, spațiu (antene) și putere;
- tehnicile cu spectru împrăștiat folosesc o lărgime de bandă de câteva ordine de mărime mai mare decât cea minim necesară unui utilizator, care este folosită simultan de mai mulți utilizatori.
- această abordare este extrem de ineficientă spectral pentru un singur utilizator, dar se dovedește extrem de eficientă pentru sisteme care deservesc mai mulți utilizatori simultan și sunt afectate de interferența de acces multiplu (multiple access interference – MAI).
- împrăștierea în frecvență („spreading”) este realizată cu ajutorul unor secvențe digitale, care trebuie să aibă proprietăți speciale, mai ales în ceea ce privește ortogonalitatea lor relativă, pentru a permite separarea la recepție a transmisiilor „suprapuse”.
- această ortogonalitate relativă între semnalele diferiților utilizatori, și/sau ale diferitelor grupuri de utilizatori, poate fi realizată în mai multe dimensiuni, prin utilizarea consecutivă a două sau chiar trei astfel de secvențe de împrăștiere.
- de asemenea, datorită acestor proprietăți, secvențele de împrăștiere asigură transmisiilor în care sunt utilizate „imunități” la unele tipuri de semnale interferente.
- datorită faptului că folosesc un spectru de frecvență mai mare decât cel minim necesar, transmisiile ce utilizează această tehnică se mai numesc cu „spectru împrăștiat prin secvență directă – DS-SS”
- datorită faptului că permit utilizarea aceleiași benzi extinse de frecvență de către mai mulți utilizatori, ce folosesc secvențe de împrăștiere de aceeași lungime care sunt ortogonale între ele, aceste transmisi se numesc „cu acces multiplu prin diviziune în cod – CDMA”
- deoarece semnalele „împrăștiate” cu acest tip de secvențe ocupă întreaga lărgime de bandă alocată unui grup de utilizatori, acest tip de transmisii gestionează resurse de tip cod-timp-frecvență.

1. Secvențe de împrăștiere (Spreading Sequences)

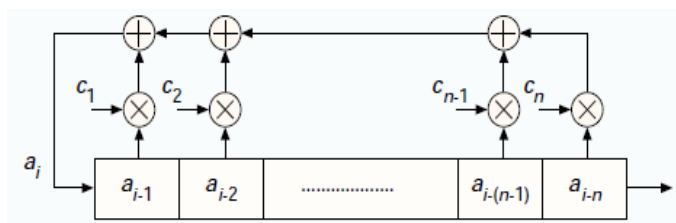
- sunt secvențe binare de lungime N^s , ale căror biți sunt numiți „chips” (de obicei în reprezentare bipolară ± 1), care sunt generate după o cîte o regulă specifică fiecărui tip de secvență;
- frecvența cu care sunt transmiși biții secvenței este $f_{ch} = N^s \cdot f_s$, f_s frecvența de simbol a transmisiei care trebuie împrăștiată, iar perioada de repetiție egală cu N^s perioade de chip, T_{ch} .

Mărimi care caracterizează secvențele de împrăștiere

- lungimea secvenței de împrăștiere, N^s , care este numărul de chipuri după care secvența se repetă
- numărul de chipuri U care se aplică secvenței de date care trebuie să fie împrăștiată
- funcția de intercorelație R_c (cross-corelation) a două astfel de secvențe se calculează înmulțind chip cu chip cele două secvențe (reprezentate în bipolar ± 1 sau -1), făcând suma produselor și împărțind-o la N^s – operații echivalente cu produsul de convoluție
- funcția de autocorelație $R_a(k)$ a unei secvențe se calculează înmulțind chip cu chip două copii ale secvenței, defazate cu k perioade de chip, făcând suma produselor și împărțind-o la N^s .
- Prezintă interes valorile lui $R_a(k)$ pentru $k = 0$ și $k \neq 0$. - nedecalată și resp. decalată

1.1 Secvențe de tip „maximal-length” (m-sequences)

- sunt secvențele cu lungime maximă ce pot fi generate de structuri cu registru de deplasare și reacție liniară (LFSR), vezi figura 1.



- mai sunt denumite și „PN codes” sau „PN-sequences” sau „m-sequences”

Figura 1. Structura generatorului unei m-secvențe

- secvența de biți a_i este generată conform formulei recurente (1) în care toți termenii sunt binari (0 sau 1):

$$a_i = \sum_{k=1}^n c_k a_{i-k} \quad (1)$$

- funcția (polinomul) generatoare a unei astfel de secvențe este de tipul (2), în care D este operatorul de întârziere cu o perioadă de chip:

$$G(D) = \frac{\sum_{k=1}^n c_k D^k (a_{-k} D^{-k} + \dots + a_{-1} D^{-1})}{1 + \sum_{k=1}^n c_k D^k} = \frac{g_0(D)}{f(D)} \quad (2)$$

- $f(D)$ este polinomul generator al circuitului cu registru de deplasare cu reacție liniară (LFSR), de care depinde vectorul conexiunilor $\{c_1, \dots, c_n\}$
- polinomul $g_0(D)$ depinde și de vectorul stării inițiale $\{a_n, \dots, a_1\}$, care determină defazajul inițial al secvenței, față de o stare de referință cunoscută de toate echipamentele implicate
- secvența LFSR este periodică de perioadă $N \leq 2^n - 1$
- secvențele LFSR de tip maximal sunt cele cu perioada $N = 2^n - 1$, pentru vectorul inițial nenul.
- o condiție ca $G(D)$ să genereze o m-sequence este ca $f(D)$ să fie ireductibil (primitiv), adică să nu poată fi descompus în factori diferiți de el însuși și de 1.
- în teorie se arată că numărul polinoamelor primitive de grad n , $N_p(n)$, este dat de (2'), unde P_i , $i = 1, 2, \dots, k$, reprezintă factorii primi ai descompunerii numărului $2^n - 1$.

$$N_p(n) = \frac{2^n - 1}{n} \cdot \prod_{i=1}^k \frac{P_i - 1}{P_i}; \quad (2')$$

- *funcția de autocorelație* indică gradul de corespondență între o secvență și o copie a ei deplasată cu k perioade de chip.
- valoarea funcției de autocorelație pentru aceste secvențe se obține cu relația (3), în care a'_n reprezintă valoarea bitului în reprezentare bipolară, iar $\delta(k)$ simbolul lui Kronecker:

$$R_a(k) = \frac{1}{N^s} \sum_{n=1}^{N^s} a'_n a'_{n+k} = \delta(k); \quad a'_n = 1 - 2a_n; \quad a_n \in \{0; 1\}; \quad (3)$$

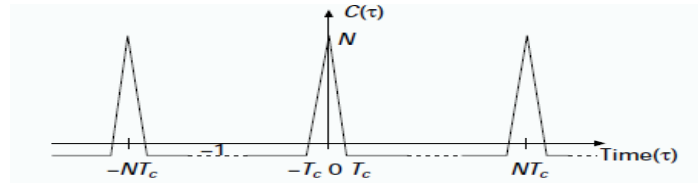
- trecând în timp continuu, dacă $p(t)$ este forma de undă corespunzătoare reprezentării bipolare a chipurilor de perioadă T_c , și definim funcția:

$$q(\tau) = \begin{cases} 1 - \frac{|\tau|}{T_c}; & |\tau| \leq T_c \\ 0; & \text{in rest} \end{cases} \quad (4)$$

- funcția de autocorelație $C(\tau)$ poate fi exprimată, pentru $N^s \gg 1$, de expresia (5), având forma prezentată în figura (2).

$$C(\tau) = N^s \sum_i q(\tau - iT_c); \quad (5)$$

Figura 2. Funcția de autocorelație a lui $p(t)$



- proprietățile funcției de autocorelație sunt utilizate pentru a alege secvențe care pot fi deosebite de secvențe pur aleatoare generate de zgomote;
- de asemenea, funcția de autocorelație evită sincronizările false ale unei secvențe la recepție.
- valoarea lui $R_a(0) = N^s$ sau $R_a(0) = 1$ (după normarea cu N^s), iar $R_a(k) = -1$, sau $R_a(k) = -1/N^s$, (după normarea cu N^s), pentru $k \neq 0$.
- *funcția de intercorelație* a două secvențe de cod este o măsură a coincidenței între două secvențe diferite a'_n și b'_n și se calculează cu relația:

$$R_c(k) = \frac{1}{N^s} \sum_{n=1}^{N^s} a'_n b'_{n+k}; \quad (6)$$

- valoarea funcției de intercorelație între oricare pereche de două secvențe maximale de lungime N^s , dintr-un set de M secvențe, este mărginită superior conform relației (7):

$$|R_c(k)| < \frac{M-1}{MN^s-1} \cong \frac{1}{N^s} \quad (7)$$

- secvențele de tip **m**, cu offseturi (condiții inițiale) diferite pot fi utilizate pentru identificarea stațiilor de bază și a celor mobile atât în downlink cât și în uplink.
- alte tipuri de secvențe de împrăștiere utilizate (sau luate în considerare pentru a fi utilizate) în

transmisiile DS-SS sunt secvențele Gold și secvențele Kasami; secvențele Gold se obțin prin combinarea unei secvențe maximale a , de lungime N , cu o variantă a sa a' , obținută prin decimarea de q ori a secvenței a și repetarea biților obținuți, unde q și N sunt relativ prime.

- aceste secvențe prezintă trei valori ale funcției de intercorelație și permit generarea unui număr foarte mare de secvențe distincte.

1.2. Secvențe de tip Walsh-Hadamard

- secvențele de tip Walsh-Hadamard (WH) fac parte din categoria secvențelor care pot genera funcții ortogonale.

- funcțiile ortogonale au proprietatea (8), în care $\varphi_i(kT_c)$ și $\varphi_j(kT_c)$ sunt membrii al i -lea și al j -lea ai unui set de funcții ortogonale, N^s e lungimea funcțiilor din set, iar T_c e durata unui chip:

$$\sum_{k=0}^{N^s-1} \varphi_i(kT_c) \cdot \varphi_j(kT_c) = 0; \quad i \neq j; \quad (8)$$

- funcțiile WH sunt reprezentate de liniile unor matrici speciale, numite matrici Hadamard, în reprezentare unipolară.

- Aceste matrici conțin o linie care are numai elemente de 0, iar restul liniilor au un număr egal de „0” și de „1”.

- funcțiile Hadamard pot fi construite după procedura descrisă de relația (9.a) în care bara superioară semnifică negarea biților care compun submatricea respectivă.

$$H_{2N^s} = \begin{bmatrix} H_{N^s} & H_{N^s} \\ H_{N^s} & \bar{H}_{N^s} \end{bmatrix}; \text{ a. } H_1 = [0]; \quad H_2 = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}; \quad H_4 = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 0 \end{bmatrix}; \text{ b.} \quad (9)$$

- funcțiile WH (în reprezentare bipolară ± 1) au funcția de intercorelație nulă, deci sunt perfect ortogonale dacă sunt sincronizate cu eroare față mai mică de fracțiuni din perioada unui chip

- în cazul în care există eroare de sincronizare mare între secvența de la recepție și cea de la emisie, corelația dintre ele poate da valori de 1, egală cu cea a funcției de autocorelație pentru sincronizare quasi-perfectă, deoarece unele secvențe se obțin prin shiftarea altor secvențe cu un număr întreg de perioade de chip,

- funcția de autocorelație $R_a(0) = 1$,

- în transmițătorul DS-SS al BS (DL) fiecare simbol de date este „împrăștiat” de o astfel de funcție (sau secvență) WH, iar factorul de împrăștiere este egal cu N^s .

- datorită proprietăților de mai sus, aceste coduri de împrăștiere ortogonale pot fi folosite doar dacă toate emițătoarele și receptoarele implicate sunt sincronizate în timp cu o eroare mai mică decât fracțiuni din perioada de chip, fapt posibil în downlink (synchronous CDMA), unde mai multe receptoare (UT) se sincronizează cu același emițător (BS).

- în uplink, BS nu poate sincroniza tactul local, cu frecvența nominală f_{chip} , cu semnalele de tact cu frecvența nominală f_{chip}^t folosite în procesul de împrăștiere de cele T UT-uri care sunt conectate; de aceea secvența locală din BS poate fi decalată cu mai mult de o perioadă de chip față de cele folosite de emițătoare

- dacă secvența locală folosită la „de-împrăștiere” este decalată cu o perioadă de chip sau mai mult (fenomen ce apare în uplink), ea va da un produs de corelație nenul cu mesajul codat cu altă secvență (a altui utilizator), și nu cu mesajul împrăștiat la emisie cu varianta ei nedecalată (a UT pe care vrem să-l recepționăm), ceea ce va conduce la obținerea unui alt mesaj la recepție; se va reveni;

- din cele de mai sus rezultă că procesele de împrăștiere care realizează extinderea (parțială) a benzii de frecvență și procesul de individualizare a utilizatorului, sau a grupului de utilizatori, nu pot fi și nu sunt întotdeauna realizate de aceeași secvență; și asupra acestui aspect se va reveni.

- de aceea în uplink (UL), funcțiile WH pot fi folosite într-o altă manieră, prin generarea de simboluri de modulare ortogonale (orthogonal modulation symbols), a căror probabilitate de eroare este mai puțin afectată de (sunt mai „rezistente” la) erorile de sincronizare

- fluxul de biți modulatori este împărțit în grupe de câte n biți, generând un simbol non-binar, căruia i se asociază una dintre cele $N^s = 2^n$ funcții Walsh, care se transmite ca o succesiune de N^s chipuri în locul grupei de n biți, realizându-se o împrăștiere cu factor de N^s/n .

- demodularea fiecărui simbol poate fi realizată prin filtrarea cu N^s filtre adaptate (care efectuează de fapt un produs de corelație)

- deoarece toți utilizatorii folosesc același set de simboluri, semnalele emise de aceștia nu pot fi separate la recepție. Pot fi separate doar simbolurile între ele, dar nu și sursele lor (UT-urile).
- de aceea în uplink semnalul fiecărui utilizator mai este împrăștiat cu o secvență PN (m-sequence), care are rolul de a duce factorul de împrăștiere la N^s , realizând o împrăștiere suplimentară cu un factor de n . Această a doua împrăștiere realizează o (pseudo)ortogonalizare între utilizatori, deoarece secvențele PN sunt pseudo-ortogonale.
- ansamblul format din generarea de simboluri modulatorie ortogonale și împrăștierea suplimentară cu secvențe PN este cunoscut sub numele de asynchronous CDMA
- secvențele PN sunt (aproximativ) necorelate statistic iar sumarea la recepție a unui număr mare de secvențe PN generează interferența de acces multiplu (*Multiple Access Interference-MAI*), care poate fi aproximată cu un zgomot gaussian, conform teoremei limitei centrale, fiind suma unor variabile aleatoare cu aceeași distribuție statistică. dar (quasi-)independente.
- dacă toți utilizatorii sunt recepționați cu același nivel al puterii, atunci pătratul dispersiei semnalului aleator format prin MAI (*Multiple Access Interference*), care e proporțional cu puterea pentru semnale aleatoare, crește direct proporțional cu numărul acestora.
- de aceea în UL, pentru CDMA asincronă, semnalele celorlalți utilizatori vor apărea ca “interferențe” pentru semnalul utilizatorului dorit, iar nivelul MAI va fi proporțional cu numărul de utilizatori autorizați care utilizează banda de frecvență respectivă.

2. Principiul CDMA (DS-SS)

- funcțiile care trebuie realizate de către o transmisie CDMA într-un sistem celular sunt: împrăștierea în frecvență, separarea (prin ortogonalitatea în cod) a utilizatorilor dintr-o celulă (sector), separarea între celule (sau sectoare sau purtătoare de canal, adică grupuri de utilizatori) și separarea sensurilor de transmisie (duplexing).
- din considerente didactice, principiul DS-SS va fi explicat considerând legătura de uplink a unei singure celule (sector sau purtător de canal), în care o aceeași secvență de împrăștiere asigură atât (pseudo)ortogonalitatea utilizatorilor, cât și împrăștierea în frecvență.
- diferențele între modalitățile de asigurare a împrăștierei în frecvență și pseudo-ortogonalității pentru cele două sensuri de transmisie, precum și modalitățile de separare a celulei (sectorului) și sensului de transmisie vor fi specificate ulterior.

2.1. Tehnica cu spectru împrăștiat prin secvența directă (*Direct Sequence Spread Spectrum DS-SS*)

- această tehnică împrăștie spectrul unui semnal A+PSK bandă de bază prin următoarele operații:

- înmulțirea coordonatelor I și Q cu secvența de împrăștiere, cu frecvența chipurilor:

$$f_c = N^s \cdot f_s \quad (10)$$

- filtrarea RRC a semnalelor $I_s(kT_c)$ și $Q_s(kT_c)$ astfel obținute;
- modularea semnalelor $I_s(t)$ și $Q_s(t)$ astfel obținute pe purtătoarele $\sin\omega_p t$ și $\cos\omega_p t$.

- ecuația semnalului modulat astfel obținut este:

$$s_{ss}(t) = A_k(t)p_t(t)\cos(\omega_i t + \Phi_k) \quad (11)$$

- filtru RRC-TJ are frecvența de tăiere $f_t = f_c(1+\alpha)/2$
- semnalul modulat are modulul nivelului maxim egal cu cel al semnalului modulator
- schema bloc de principiu a modulatorului DS-SS este prezentată în figura 3.

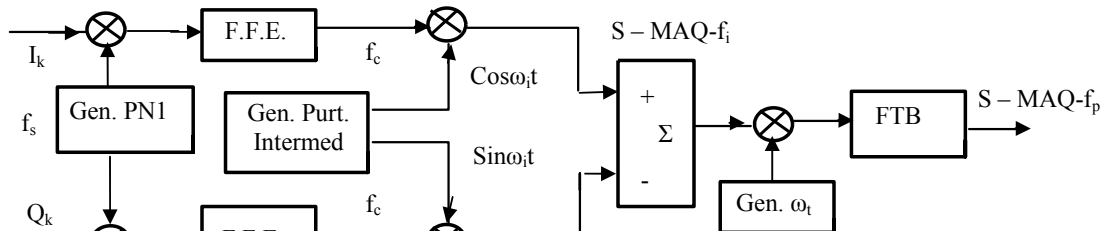


Figura 3. Schema bloc a modulatorului DS-SS

- pe sensul uplink se folosește OQPSK (se defazează tactul de simbol folosit la modularea axei Q cu o semiperioadă de chip față de cel folosit la modularea axei I) pentru a scădea valoarea PAPR și a reduce astfel distorsiunile introduse de amplificatorul final al UT, vezi cursul de TM.
- unele variante folosesc secvențe de împrăștiere diferite pe axele I și Q

2.2. Spectrul semnalului modulat DS-SS

- semnalul modulator BB are spectrul de tip sinus atenuat specific unei modulații QAM, vezi cursul

de TM, având lobul spectral principal cuprins între $-f_s$ și f_s ,

- în urma înmulțirii nivelelor I și Q cu $p(kT_c)$ spectrul semnalului modulator rezultat are lobul principal cuprins între $-f_c$ și f_c , iar lărgimea de bandă a semnalului modulat DS-SS pe purtătoarea de canal este W_{ss} .

- factorul de împrăștiere al benzii, PG (numeric egal cu câștigul procesării PG, vezi considerațiile de la demodulare), este:

$$PG = \frac{f_c \cdot (1 + \alpha)}{f_s \cdot (1 + \alpha)} = \frac{T_s}{T_c} = \frac{W_{ss}}{W_{QAM}} = N^s \quad (12)$$

2.3 Demodularea semnalului DS-SS

- principiul demodulării este reprezentat schematic în figura 4

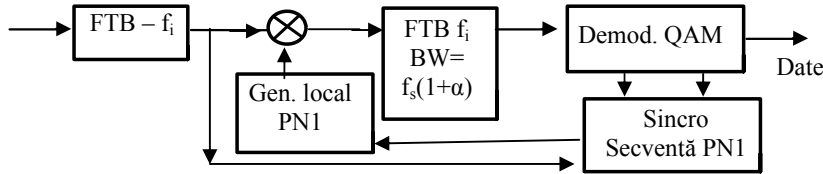


Figura 4. Schema bloc a demodulatorului DS-SS

- demodularea semnalului DS-SS implică mai întâi translatarea acestuia pe frecvență intermediară și o filtrare trece-bandă cu un filtru de bandă largă (mai mare ca $f_c(1+\alpha)$).

- apoi semnalul recepționat este înmulțit cu secvența de împrăștiere $p'(t)$ generată local. Presupunând că aceasta este **corect sincronizată**, înmulțirea conduce la:

$$s_s(t) = A_k(t)p_t(t)p'_t(t)\cos(\omega_i t + \Phi_k) \quad (13)$$

- ținând cont de faptul că în urma medierii asigurate de un filtru TB, cu lărgimea de bandă $f_s(1+\alpha)$ mult mai mică decât f_c , aflat la intrarea demodulatorului QAM, putem presupune că p_t^2 mediat = 1, semnalul filtrat va fi semnalul QAM $A_k\cos(\omega_i t + \Phi)$ axat pe frecvența intermediară având lărgime de bandă $f_s(1+\alpha)$.

- operația de înmulțire cu secvența locală de împrăștiere sincronizată este numită și „despreading”.

- ca urmare a acestei operații lărgimea de bandă a semnalului se reduce de la W_{ss} la W_{QAM} așa cum se arată în figura 5.

- dacă însă secvența locală folosită la „despreading” nu este sincronizată cu secvența folosită la emisie pentru „spreading” atunci semnalul rezultat are o amplitudine foarte mică (de $1/N$ ori mai mică), vezi proprietățile secvențelor SS, ceea ce face ca acesta să genereze o probabilitate de eroare foarte mare după demodularea QAM, datorită valorii reduse a raportului semnal/(zgomot+interferență), SINR.

- în cazul împrăștierii cu secvențe diferite pe axele I și Q, atunci deîmprăștierea se face după demodularea QAM, blocul de deîmprăștiere fiind plasat după blocul de sondare, operația fiind efectuată în banda de bază, separat pe cele două axe

2.3 Reducerea puterii semnalelor interferente

- un efect important al operațiilor de „spreading-despreading” este reducerea semnificativă a raportului între puterea semnalului util și puterea unui semnal aditiv interferent $i(t)$ de bandă îngustă (comparabilă cu a semnalului QAM neîmprăștiat) ce afectează semnalul DS-SS recepționat la trecerea prin canalul de transmisie:

$$s_r(t) = s_{ss}(t) + i(t) \quad (14)$$

- în urma operației de despreading puterea semnalului util (care era distribuită într-o bandă largă W_{ss}) este concentrată într-o bandă îngustă W_{QAM} (banda semnalului util înainte de împrăștiere), iar puterea semnalului interferent (care era distribuită într-o bandă îngustă B_{interf}) este „împrăștiată” într-o bandă largă $PG \cdot B_{interf}$. – vezi figura 5 și relația (15).

$$s_d(t) = s_{ss}(t) \cdot p_t(t) + i(t) \cdot p_t(t) = A_k \cos(\omega_i t + \Phi_k) + i(t) \cdot p_t(t) \quad (15)$$

- în urma trecerii prin filtrul TB de la intrarea demodulatorului QAM precum și a filtrării TJ efectuate de demodulatorul QAM ($f_i = f_s(1+\alpha)$), semnalul util își păstrează puterea de după despreading în timp ce puterea semnalului interferent scade de PG ori. Acest fenomen se datorează faptului că semnalul util a trecut prin ambele procese de „spreading-despreading”, pe când semnalul interferent a trecut doar prin operația de împrăștiere”.

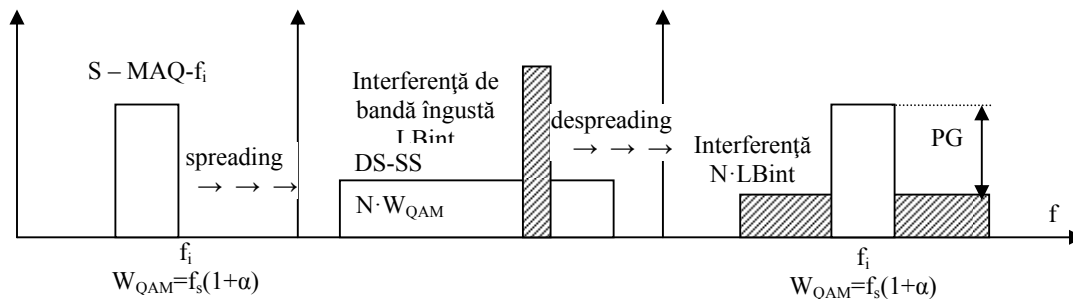


Figura 5 Reducerea puterii interferenței de bandă îngustă - reprezentare schematică

- factorul de scădere a puterii semnalelor interferente în banda utilă se numește „processing gain” și este dat de raportul benzilor semnalului „împrăștiat” și celui original. Acest factor este cu atât mai mare cu cât factorul de împrăștiere este mai mare.
- această comportare este aplicabilă și semnalelor interferente de bandă largă (comparabilă cu lărgimea de bandă a semnalului împrăștiat) exterioare celei, deoarece ele sunt necorelate cu secvența de împrăștiere.
- concluzionând putem spune că procesul de spreading-despreading conduce la o creștere a raportului semnal/interferențe de PG ori (sau a valorii SIR cu $10 \lg PG$) la ieșirea din demodulatorul QAM înainte de blocul de decizie.
- studiile prezentate în literatură arată că această creștere nu apare în cazul zgomotului gaussian. În cazul acestui zgomot valoarea SNR nu se modifică față de o transmisie ce nu folosește tehnica DS-SS pe același canal.
- acest fenomen se poate explica principial prin faptul că zgomotul gaussian este un fenomen aleator cu proprietăți statistice și spectrale apropiate de cele ale secvențelor de „împrăștiere” utilizate pentru semnalul util; de aceea, înmulțirea semnalului de zgomot gaussian la recepție cu secvența de împrăștiere nu mai conduce la o mărire a benzii de frecvență a acestuia, și implicit la scăderea densității spectrale de putere N_0 .